

EDUARDO CAMPANA BARBOSA

INFERÊNCIA VIA *BOOTSTRAP* NA *CONJOINT ANALYSIS*

Tese apresentada à Universidade Federal de Viçosa, como parte das exigências do Programa de Pós-Graduação em Estatística Aplicada e Biometria, para obtenção do título de *Doctor Scientiae*.

VIÇOSA
MINAS GERAIS – BRASIL
2017

**Ficha catalográfica preparada pela Biblioteca Central da Universidade
Federal de Viçosa - Câmpus Viçosa**

T

B238i
2017 Barbosa, Eduardo Campana, 1988-
 Inferência via *Bootstrap* na *Conjoint Analysis* / Eduardo
Campana Barbosa. – Viçosa, MG, 2017.
 x, 57f. : il. (algumas color.) ; 29 cm.

Inclui apêndices.

Orientador: Carlos Henrique Osório Silva.

Tese (doutorado) - Universidade Federal de Viçosa.

Inclui bibliografia.

1. Estatística - Modelos matemáticos. 2. Marketing.
3. Método *Bootstrap*. I. Universidade Federal de Viçosa.
Departamento de Estatística. Programa de Pós-graduação em
Estatística Aplicada e Biometria. II. Título.

CDD 22. ed. 519.5

EDUARDO CAMPANA BARBOSA

INFERÊNCIA VIA *BOOTSTRAP* NA *CONJOINT ANALYSIS*

Tese apresentada à Universidade Federal de Viçosa, como parte das exigências do Programa de Pós-Graduação em Estatística Aplicada e Biometria, para obtenção do título de *Doctor Scientiae*.

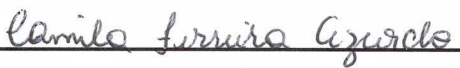
APROVADA: 14 de dezembro de 2017.



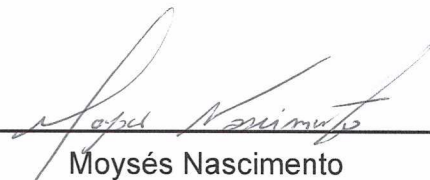
Valéria Paula Rodrigues Minim



João de Deus Souza Carneiro



Camila Ferreira Azevedo
Coorientadora



Moysés Nascimento
Coorientador



Carlos Henrique Osório Silva
Orientador

Aos meus pais Jorge Luiz Resende Barbosa e Denise Gama
Campana Barbosa. A minha irmã Gabriela Campana Barbosa.

Agradecimentos

- Em primeiro lugar a Deus, por permitir que mais uma importante etapa de minha vida possa ser concluída.
- A meu pai (Jorge) e minha mãe (Denise), que desde a época de graduação, com muito esforço, permitiram com que eu pudesse me dedicar inteiramente aos estudos, o que influenciou meu interesse profissional pela área acadêmica e consequentemente para o ingresso no programa de pós-graduação em Estatística Aplicada e Biometria da UFV. Agradeço também a minha irmã (Gabriela), que sempre esteve ao meu lado, mesmo à distância, todas as vezes que precisei. Essa conquista é nossa.
- Ao meu orientador, professor Carlos Henrique Osório Silva (C.H.O.S.), por sua paciência, dedicação e amizade. Fica o meu agradecimento por todos os conhecimentos transmitidos e pelas diversas vezes que me ajudou com algum problema profissional/pessoal, sempre com boa vontade e ótimas sugestões.
- Aos meus co-orientadores, Camila Ferreira Azevedo e Moysés Nascimento, pela atenção e pelas valiosas contribuições, que muito melhoraram o meu trabalho e também minha percepção sobre pesquisa científica. Adicionalmente, agradeço-os pela amizade e ajuda desde que ingressei no programa de Pós-graduação em Estatística Aplicada e Biometria aqui da UFV.
- A todos os professores e funcionários do Departamento de Estatística Aplicada e Biometria da Universidade Federal de Viçosa (DET-UFV), pelo acolhimento e ensinamentos. Em especial a melhor secretária do DET - UFV, Maria Anita Paiva, que me ajuda a preencher os papéis de estágio probatório, processo de treinamento, bancas de monitoria etc., mostrando-se sempre atenciosa, eficiente e com educação e “paciência”.
- Aos meus poucos, mas verdadeiros amigos, que eu sei que sempre poderei contar. Obrigado por me motivarem. Fernando (Abelha), Filipe (Formiga), Heider (Doquinha 1), em memória de Hugo (Papai), Lucas (Banco do Brasil),

Mateus (exemplo de profissional), Osvaldo (meu ídolo), Ricardo (Zé B.), Romulo (Manuli), Ronaldo (Perigoso) e Sérgio (Doquinha 2).

Sumário

Lista de Figuras.....	vii
Lista de Tabelas	viii
Resumo	ix
Abstract.....	x
 INTRODUÇÃO GERAL.....	 1
 REVISÃO DE LITERATURA.....	 3
1. <i>Conjoint Analysis</i>	3
2. <i>Bootstrap</i>	7
REFERÊNCIAS BIBLIOGRÁFICAS.....	9
 CAPÍTULO I	 14
INFERÊNCIA VIA <i>BOOTSTRAP</i> NA <i>RATINGS-BASED CONJOINT ANALYSIS</i>	14
Resumo:	14
Abstract.....	14
1. Introdução	14
2. Material e métodos.....	16
2.1 Dados experimentais analisados	16
2.2 Ratings-Based Conjoint Analysis (RBCA).....	18
2.3 Princípios básicos do <i>Bootstrap</i> não paramétrico	19
2.4 Método <i>Residuals Bootstrap</i> na RBCA	20
2.5 Aspectos computacionais.....	22
3. Resultados e discussão	22
4. Conclusões	27
5. Referências.....	27
Apêndice I.....	31
 CAPÍTULO II	 35
INFERÊNCIA VIA <i>BOOTSTRAP</i> NA <i>CHOICE-BASED CONJOINT ANALYSIS</i>	35
Resumo	35
Abstract.....	35
1. Introdução	35
2. Material e Métodos	37

2.1 Dados Experimentais Analisados.....	37
2.2 Choice-Based Conjoint Analysis (CBCA).....	38
2.3 Princípios Básicos do <i>Bootstrap</i> não paramétrico	40
2.4 Método <i>Paired Bootstrap</i> na CBCA.....	41
2.5 Aspectos Computacionais.....	42
3. Resultados e Discussão	42
4. Conclusões.....	48
5. Referências Bibliográficas	49

Lista de Figuras

CAPÍTULO I

Figura 1: Tratamentos. (1) Iogurtes, (2) Bebidas achocolatadas, (3) Queijos pratos ilustrados..17

Figura 2: Histograma das distribuições empíricas dos estimadores do modelo da RBCA.....25

Figura 3: Histograma das distribuições empíricas das importâncias relativas.....26

CAPÍTULO II

Figura 1: Histograma das distribuições empíricas dos estimadores do modelo da CBCA.....44

Figura 2: Histograma das distribuições empíricas das probabilidades de escolha.....46

Figura 3: Histograma das distribuições empíricas da razão de escolha dos atributos.....48

Lista de Tabelas

REVISÃO DE LEITERATURA

Tabela 1: Atributos e níveis do tratamento analisado (camisa).....	3
Tabela 2: Esquema Fatorial Completo para o exemplo da camisa.....	3

CAPÍTULO I

Tabela 1: Delineamento dos Tratamentos em estudo.....	17
Tabela 2: Estimativas e estatísticas fornecidas pelo método <i>Residuals Bootstrap</i>	22
Tabela 3: Intervalos baseados na distribuição t-Student e intervalos Percentis <i>Bootstrap</i>	24
Tabela 4: Estimativas e estatísticas <i>Bootstrap</i> para as importâncias relativas dos atributos.....	25

CAPÍTULO II

Tabela 1: Delineamento dos Tratamentos em estudo.....	38
Tabela 2: Estimativas e estatísticas fornecidas pelo método <i>Paired Bootstrap</i>	42
Tabela 3: Intervalos baseados na distribuição Normal Padrão e intervalos Percentis <i>Bootstrap</i>	43
Tabela 4: Estimativas e estatísticas <i>Bootstrap</i> para as probabilidades de escolha.....	44
Tabela 5: Estimativas e estatísticas <i>Bootstrap</i> para a razão de escolha dos atributos.....	46

Resumo

BARBOSA, Eduardo Campana, D. Sc., Universidade Federal de Viçosa, dezembro de 2017. **Inferência via *Bootstrap* na *Conjoint Analysis***. Orientador: Carlos Henrique Osório Silva. Coorientadores: Moysés Nascimento e Camila Ferreira Azevedo.

A presente tese teve como objetivo introduzir o método de reamostragem com reposição ou *Bootstrap* na *Conjoint Analysis*. Apresenta-se no texto uma revisão conceitual (Revisão de Literatura) sobre a referida metodologia (*Conjoint Analysis*) e também sobre o método proposto (*Bootstrap*). Adicionalmente, no Capítulo I e II, define-se a parte teórica e metodológica da *Conjoint Analysis* e do método *Bootstrap*, ilustrando o funcionamento conjunto dessas abordagens via aplicação real, com dados da área de tecnologia de alimentos. Inferências adicionais que até então não eram fornecidas no contexto clássico ou frequentista podem agora ser obtidas via análise das distribuições empíricas dos estimadores das Importâncias Relativas (abordagem por notas) e das Probabilidades e Razão de Escolhas (abordagem por escolhas). De forma geral, os resultados demonstraram que o método *Bootstrap* forneceu estimativas pontuais mais precisas e tornou ambas as abordagens da *Conjoint Analysis* mais informativas, uma vez que medidas de erro padrão e, principalmente, intervalos de confiança puderam ser facilmente obtidos para certas quantidades de interesse, possibilitando a realização de testes ou comparações estatísticas sobre as mesmas.

Abstract

BARBOSA, Eduardo Campana, D. Sc., Universidade Federal de Viçosa, December, 2017. **Inference by Bootstrap in Conjoint Analysis**. Advisor: Carlos Henrique Osório Silva. Co-Advisors: Moysés Nascimento and Camila Ferreira Azevedo.

The aim of this thesis was introduce the Bootstrap resampling method in Conjoint Analysis. We present in the text a conceptual review (Literature Review) about this methodology (Conjoint Analysis) and also about the proposed method (Bootstrap). In addition, in Chapter I and II, the theoretical and methodological aspects of Conjoint Analysis and the Bootstrap method are defined, illustrating the joint operation of these approaches via real application, with data from the food technology area.. Additional inferences have not been provided in the classic or frequentist context can now be obtained by analyzing the empirical distributions of Relative Importance (ratings based approach) and Probability and Choice Ratio (choice based approach) estimators. Overall, the results demonstrated that the Bootstrap method provided more accurate point estimates and made both Conjoint Analysis approaches more informative, since standard error measures, and mainly confidence intervals, could be easily obtained for certain quantities of interest, making it possible to perform statistical tests or comparisons on them.

INTRODUÇÃO GERAL

Diversos fatores contribuem para que a competitividade entre as empresas aumente, dentre estes a evolução tecnológica, a mudança de percepção dos consumidores em termos de qualidade de produtos, serviços, conceitos (dentre outros aos quais se denomina como um tratamento), além da busca por satisfação pessoal. Nota-se ainda que o nível de exigência dos consumidores em relação aos tratamentos aumenta continuamente, fato que reflete diretamente a necessidade que as empresas possuem de estabelecer novas estratégias competitivas e mercadológicas para fidelizar os seus clientes.

Estudos são conduzidos na área de *marketing* no intuito de tentar entender a valiosa relação entre empresa e os desejos dos consumidores por alguns tratamentos. Segundo Kotler e Armstrong (2008), a lucratividade e o sucesso de uma empresa são objetivos específicos que estão associados ao bom relacionamento e a uma comunicação eficiente com o seu cliente. Isso pode ser explicado devido o consumidor ser considerado como o destino final do tratamento, tornando evidente o quão importante e forte é a influência deste para que uma empresa tenha sucesso perante a agressividade do mercado.

Sabe-se ainda que as empresas precisam oferecer a seus clientes (consumidores) o que estes realmente desejam, com qualidade e eficiência, visando à sua satisfação (do cliente) e também maior rentabilidade (da empresa). No entanto, a principal pergunta a ser feita é: que tipo de tratamento o cliente realmente deseja? Logo, o primeiro ponto para responder esta questão consiste em entender os desejos e consequentemente as preferências dos clientes. É neste contexto que a técnica estatística denominada *Conjoint Analysis* (Análise Conjunta de Fatores) mostra-se uma ferramenta de fundamental importância para o cenário empresarial, seja para auxiliar em tomadas de decisões ou para o desenvolvimento de novos tratamentos.

A presente tese encontra-se estruturada da seguinte forma: Na próxima seção será apresentada uma Revisão de Literatura sobre a técnica *Conjoint Analysis* e também sobre o método *Bootstrap* (Efron, 1979). Adicionalmente, apresenta-se uma aplicação prática da referida metodologia, baseada em Notas (Capítulo I) e baseada em Escolhas (Capítulo II), em conjunto com o método *Bootstrap*. A introdução desse método na *Conjoint Analysis* permitirá o acesso da distribuição empírica de qualquer estatística de

interesse, como as Importâncias Relativas (para abordagem baseada em Notas), as Probabilidades e as Razões de Escolha (para a abordagem baseada em Escolhas). Esse tipo de informação, até então não disponibilizada no enfoque Frequentista ou Clássico, permitirá a realização de inferências (testes de hipóteses e comparações estatísticas) para tais quantidades, o que é um resultado de grande interesse, principalmente sobre o ponto de vista prático.

REVISÃO DE LITERATURA

1. Conjoint Analysis

A *Conjoint Analysis* é uma técnica estatística que permite avaliar a intenção de compra dos consumidores em relação a diversos tratamentos (produtos, serviços, conceitos etc.). Esses tratamentos são formados pela combinação de características específicas (níveis dos atributos), e o objetivo principal da referida técnica é auxiliar o pesquisador a entender a importância de cada uma destas características para o consumidor, no momento da compra ou da aquisição de um tratamento (Moskowitz et al., 2004). Como um exemplo ilustrativo, suponha o interesse em avaliar a influência de $r = 3$ atributos na intenção de compra por um tipo de camisa. Admita que cada um dos atributos em estudo possua 2 níveis, isto é, $m_1 = m_2 = m_3 = 2$, conforme ilustrado na Tabela 1.

Tabela 1 – Atributos e níveis do tratamento analisado (camisa).

Nível/Atributo	1	2	3
	Cor	Ilustração	Preço
1	Preta	Sim	15,00
2	Branca	Não	20,00

O total de tratamentos formado pela combinação dos m_s níveis dos $r = 1, 2, 3$ atributos é denominado como esquema fatorial completo e é definido por $J = \prod_{s=1}^r m_s$, neste caso, $J = 2^3 = 8$, conforme apresentado na Tabela 2.

Tabela 2 – Esquema Fatorial Completo para o exemplo da camisa.

Tratamento	Cor	Ilustração	Preço
1	Preto	Sim	15,00
2	Branco	Sim	15,00
3	Preto	Não	15,00
4	Branco	Não	15,00
5	Preto	Sim	20,00
6	Branco	Sim	20,00
7	Preto	Não	20,00
8	Branco	Não	20,00

Note que o aumento do número de atributos e/ou de seus respectivos níveis ocasiona o aumento do número de tratamentos, sendo este um ponto que se deve atentar

ao utilizar a metodologia *Conjoint Analysis*, uma vez que esse fato pode tornar a avaliação dos consumidores complexa e cansativa, o que diminuirá a precisão dos resultados e inferências. Hair Júnior et al. (1995) sugerem que a escolha dos atributos e níveis seja realizada com relevância e que as características escolhidas para compor os tratamentos realmente influenciem na opção de escolha do consumidor. Algumas alternativas são as seções do tipo *Focus Group* (Deliza, 1996) e o método da Soma Constante (Dias, 2010), metodologias utilizadas para auxiliar o pesquisador a definir os atributos e os níveis antes da aplicação do estudo ou os esquemas fatoriais fracionados (delineamentos de tratamentos) (Kuhfeld et al., 1997), empregados para reduzir o número de tratamentos após a definição do conjunto de alternativas.

Ressalta-se que o termo *Conjoint* (de Conjunta) consiste em estudar conjuntamente as características dos tratamentos, pois suas importâncias relativas podem ser mensuradas com uma intensidade que não é possível ser obtida quando estas são avaliadas separadamente (Churchill Jr. e Nielsen Jr., 1996). Esta é uma das principais características da *Conjoint Analysis* e, portanto, em situações práticas é adequado que os consumidores entendam que estão avaliando um tratamento e não as características específicas que os compõe separadamente.

Adicionalmente, é preciso definir como cada tratamento será apresentado aos consumidores. Segundo Silva e Bastos (2010) na realização do estudo deve escolher a forma mais adequada, seja por meio de instruções, questionários, objetos reais, bens de consumo, protótipos de produtos ou fotografias. Além disso, a avaliação deverá ocorrer em um local que seja tranquilo e confortável, com um intervalo de tempo adequado entre a apresentação de cada tratamento aos consumidores.

A ordem de apresentação dos tratamentos é também de extrema importância, uma vez que existem implícitos os efeitos da ordem de apresentação e da influência de uma amostra na avaliação subsequente (Della Lúcia et al., 2007), o que é explicado pela “psicologia do consumidor” e, associada com a expectativa dos avaliadores a cada um dos tratamentos. Uma alternativa adotada para minimizar estes problemas é apresentar os tratamentos conforme os delineamentos experimentais proposto por MacFIE et al. (1989). Nestes delineamentos a ordem de apresentação de cada tratamento é diferente para cada consumidor. Além disso, cada tratamento aparece precedido e sucedido por todos os outros tratamentos o mesmo número de vezes.

Por fim, a *Conjoint Analysis* pode ser dividida em duas técnicas de análise. Na abordagem inicial, denominada como *Ratings-Based Conjoint Analysis* (Análise

Conjunta Baseada em Notas, referida aqui como RBCA), em que os consumidores avaliam um determinado número de tratamentos e atribuem a cada um destes uma nota discreta ou contínua, dentro de um determinado intervalo. A análise estatística para este caso consiste na utilização de modelos usuais de Regressão Linear Múltipla. Essa técnica foi inicialmente introduzida na área de *Marketing* por Green e Rao (1971). Posteriormente, inúmeros trabalhos foram propostos com o intuito de: i) desenvolver novos produtos (Sands e Warwick, 1981; Green e Krieger, 1987; Yoo e Ohta, 1995), ii) avaliar e otimizar produtos já existentes (Green e Wind, 1975; Green e Srinivasan, 1978), iii) estabelecer novas estratégias competitivas [políticas de preço (Mahajan, Green e Goldberg, 1982), propaganda (Tscheulin e Helmig, 1998) e distribuição (Green e Savitz, 1994), entre outros (Gustafsson, Herrmann, e Huber, 2001). Em estudos recentes a aplicação da RBCA foi generalizada a diversas áreas do conhecimento, como a área médica (Stockwell et al., 2011), administração (Rhee e Yang, 2015) e tecnologia de alimentos (Andrade et al., 2016), entre outros.

Já a segunda abordagem da *Conjoint Analysis*, denominada como *Choice-Based Conjoint Analysis* (Análise Conjunta de Fatores Baseada em Escolhas, referida aqui como CBCA), os consumidores informam sua preferência por um ou mais tratamentos por meio de escolhas. Neste caso a análise estatística ocorre via modelo de regressão Logit. A origem da CBCA foi com o trabalho de Louviere e Woodworth (1983), que combinaram os conceitos experimentais da *Conjoint Analysis* tradicional (Luce e Tukey, 1964) com os modelos econométricos de escolha discreta (Train, 2009), mais especificamente, o modelo Logit Multinomial (MLM), inicialmente derivado por Luce (1959) e completamente especificado por McFadden (1974). Posteriormente, na área de *marketing* inúmeros trabalhos foram propostos, difundindo e consolidando essa metodologia (Kamakura e Srivastava, 1984; Johnson e Olberts, 1991; Koelemeijer e Oppewal, 1999; Moore, Louviere e Verma, 1999; entre outros). A CBCA também tem sido utilizada em outras áreas do conhecimento, como em pesquisas sociais (Adamowicz, Boxall e Williams, 1998; Bullock, Ellston e Chalmers, 1998), geografia (van der Waerden, Oppewal e Timmermans, 1993; Oppenwal, Timmermans e Louviere 1997) e Turismo/Transporte (Brandley e Gunn, 1990; Haider e Ewing, 1990). Em estudos recentes verificam-se ainda algumas aplicações na área de saúde (Ryan e Farrar, 2000; Zimmermann et al., 2013; e Utz et al., 2014). No Brasil essa metodologia é pouco explorada. Suas principais aplicações são na área de engenharia de alimentos (veja Della Lucia et al., 2010; Deliza et al., 2010).

Uma limitação da *Conjoint Analysis* tradicional (baseada em Notas e também em Escolhas) é que essa técnica não considera a heterogeneidade entre os consumidores, ou seja, pressupõe-se que a amostra em estudo possui, em geral, a mesma preferência. Portanto, as estatísticas e estimativas obtidas serão únicas e tenderão a representar o comportamento médio dos consumidores. Neste sentido, algumas abordagens têm sido desenvolvidas, dentre as quais destacamos os modelos com efeitos aleatórios e a aplicação simultânea da *Conjoint Analysis* com técnicas de clusterização (Rao, 2014). Adicionalmente, métodos de estimação Bayesianos, que procedem a *Conjoint Analysis* via modelos hierárquicos, denotados como *Hierarchical Bayes* (HB), conforme apresentados em Rossi et al. (2005), também são alternativas de análise. Nesse tipo de abordagem estimam-se os parâmetros dos modelos estatísticos a nível individual, isto é, estima-se um vetor de parâmetros para cada indivíduo.

2. *Bootstrap*

O método *Bootstrap* teve origem com o trabalho “*Bootstrap methods: another look at the jackknife*” desenvolvido por Efron (1979). Posteriormente, outros textos se tornaram referências clássicas sobre o tema, como os livros de Efron e Tibshirani (1993) e Davison e Hinkley (1997), intitulados respectivamente como: “*An Introduction to the Bootstrap*” e “*Bootstrap Methods and their Applications*”.

O termo *Bootstrap* faz referência à expressão “*pulling oneself up by one's bootstraps*” cuja tradução é “de que é possível emergir de um afogamento sendo puxado pelo cadarço do próprio sapato”. No enfoque estatístico, essa expressão sugere que pode-se obter propriedades válidas em grandes amostras (ou propriedades assintóticas) a partir de um número pequeno de observações. A ideia principal desse método é acessar a distribuição amostral de determinada estatística por reamostragem com reposição, a partir de uma amostra original ou mestra (Fox e Weisber, 2011). Mais especificamente, torna-se possível inferir e avaliar a incerteza sobre a estatística de interesse quando procedimentos teóricos: i) não podem ser aplicados ou ii) são analiticamente complexos ou ainda iii) quando o tamanho da amostra é relativamente pequeno (Davison e Kuonen, 2002). Neste sentido, no contexto da *Conjoint Analysis*, podem-se citar as Importâncias Relativas dos atributos e as Probabilidades de Escolha dos tratamentos, funções não lineares parâmetros de interesse e que exigem conhecimentos estatísticos específicos para a dedução analítica de sua distribuição amostral ou de medidas de precisão.

Em geral, no contexto Frequentista, o *Bootstrap* pode ser realizado de duas formas: paramétrica¹ e não-paramétrica. A primeira é adequada para situações em que o pesquisador possui alguma informação probabilística sobre a população de interesse, ou seja, as reamostras ou réplicas *Bootstrap* são geradas a partir de uma função de distribuição paramétrica ou de um modelo matemático, cujos parâmetros definem por completo a função densidade ou a função de probabilidade da população. Isto é: $Y_1, Y_2, \dots, Y_n \sim i. i. d. F(\mathbf{y})$ (Davison e Hinkley's, 1997). Já na segunda alternativa, não são feitas suposições teóricas sobre a distribuição da população, e as réplicas *Bootstrap* são obtidas por reamostragem com reposição da amostra original. Admite-se que a função de distribuição dos dados é desconhecida e estimada por sua respectiva

¹ No *Bootstrap* paramétrico admite-se que a distribuição dos dados seja conhecida, no entanto, os respectivos valores paramétricos são especificados a partir da amostra original ou mestra, que já foi observada. Já na simulação Monte Carlo, assume-se que a distribuição e também que os valores paramétricos sejam conhecidos, logo, os dados são gerados de forma aleatória e sem "conteúdo empírico", o que torna esse método mais apropriado para avaliações teóricas.

distribuição empírica, isto é, $Y_1, Y_2, \dots, Y_n \sim i.i.d. \hat{F}(\mathbf{y})$. Adicionalmente, considera-se que a amostra original é independente e identicamente distribuída, e por isso a ocorrência ou a seleção de qualquer observação está associada a uma probabilidade de $1/n$ ou n^{-1} (Efron, 1979).

O método *Bootstrap* possui limitações. Para que as inferências produzidas sejam válidas, algumas condições devem ser atendidas. Em particular, a amostra original deve ser aleatória e representativa da população, para minimizar o viés dos resultados. Caso contrário, variações deste método são necessárias. Para situações em que os dados são correlacionados, como tradicionalmente ocorre na análise de séries temporais, a reamostragem pode ser realizada, por exemplo, dentro de blocos, o que é conhecido como *Moving Block Bootstrap* (Kunsh, 1989). Problemas de heterocedasticidade ou não homogeneidade dos erros, principalmente nas análises de regressão, podem ser corrigidos com o que é conhecido como *Wild Bootstrap* (Liu, 1988). Além disso, como cada observação da amostra possui a mesma probabilidade de ser selecionada para compor uma réplica *Bootstrap*, possíveis *outliers* (valores discrepantes) precisam ser removidos ou avaliados na análise, o que pode ser feito via procedimento “*jackknife após o bootstrap*” (Efron, 1992).

Destaca-se que, além do *Bootstrap* e de suas variações, existem outros métodos de reamostragem já documentados na literatura, como o *Jackknife* e o *Cross-Validation* e os testes *Permutation* e *Randomization*, que se diferem do *Bootstrap* pela forma como a reamostragem é realizada. Para maiores detalhes consultar Efron e Tibshirani's (1993) e Davison e Hinkley's (1997). Adicionalmente, sob o enfoque Bayesiano, podem-se derivar as distribuições à posteriori para os parâmetros e qualquer função real destes, o que permite ao pesquisador realizar inferências similares às obtidas via método *Bootstrap*. Maiores detalhes podem ser consultados em Rossi et al. (2005) e Chapman e Feit (2015).

REFERÊNCIAS BIBLIOGRÁFICAS

- Adamowicz, W. P., M. Boxall, and M. Williams. (1998). Stated preference approaches for measuring passive use values: Choice experiments versus contingent valuation. *American Journal of Agricultural Economics* 80:64–75.
- Andrade, J. C., Nalério, É. S., Giongo, C., de Barcellos, M. D., Ares, G., e Deliza, R. (2016). Influence of evoked contexts on rating-based conjoint analysis: Case study with lamb meat. *Food Quality and Preference*, 53, 168-175.
- Brandley, M. A., and H. Gunn. (1990). Stated preference analysis of values of travel time in the Netherlands. *Transportation Research Record* 1285:78–89.
- Bullock, C. H., D. A. Elston, and N. A. Chalmers. (1998). An application of economic experiments to a land-use, deer hunting, and landscape change in the Scottish highlands. *Journal of Environmental Management* 52:335–51.
- Chapman, C., e Feit, E. M.. *R for marketing research and analytics*. Springer, 2015.
- Chernick, M. R.; LaBudde, R. A. (2014). An introduction to bootstrap methods with applications to R. John Wiley e Sons.
- Churchill Jr., G. A.; Nielsen Jr., A. C. (1996). *Marketing Research: Methodological Foundations*. 6th edition. Wisconsin: The Dryden Press, 602 p.
- Churchill Jr., G. A.; Nielsen Jr., A. C. *Marketing Research: Methodological Foundations*. 6th edition. Wisconsin: The Dryden Press, 1996. 602 p.
- Davison, A. C.; Hinkley, D. V. (1997). *Bootstrap Methods and their Applications*. Cambridge University Press, Cambridge.
- Davison, A. C.; Kuonen, D. (2002). An Introduction to the Bootstrap with Applications in R. *Statistical Computing and Statistical Graphics Newsletter*, 13(1), 6-11.
- Deliza, R. (1996). The effects of expectation on sensory perception and acceptance. 1996. 198p. (PhD thesis) - University of Reading, Inglaterra, 1996.
- Deliza, R., Rosenthal, A., Hedderley, D., e Jaeger, S. R. (2010). Consumer perception of irradiated fruit: a case study using choice- based conjoint analysis. *Journal of Sensory Studies*, 25(2), 184-200.

- Della Lucia, S. M., Minim, V. P. R., Silva, C. H. O., Minim, L. A., e Silva, R. C. S. N. (2010). Análise conjunta de fatores baseada em escolhas no estudo da embalagem de iogurte light sabor morango. *Brazilian Journal of Food technology*, 11-18.
- Della Lucia, S. M.; Minim, V. P. R.; Silva, C. H. O.; Minim, L. A. (2007). Fatores da embalagem de café orgânico torrado e moído na intenção de compra do consumidor. *Ciência e Tecnologia de Alimentos*, v. 27, p. 485-491.
- Dias, A (2010). Comparação entre os Métodos Soma Constante e Análise Conjunta de Fatores para Estimar a Importância Relativa. 65p. Dissertação (Mestrado em Estatística Aplicada e Biometria), Departamento de Estatística, Universidade Federal de Viçosa, Viçosa.
- Efron, B. (1979). Bootstrap methods: another look at the jackknife. *Annals of Statistics*, 7:1 – 26.
- Efron, B. (1981). Nonparametric standard errors and condence intervals (with discussion). *Canadian Journal of Statistic*, 9, 139-172.
- Efron, B. (1992). Bootstrap methods: another look at the jackknife. In *Breakthroughs in statistics* (pp. 569-593). Springer New York.
- Efron, B.; Tibshirani, R. J. (1993). *An Introduction to the Bootstrap*. Chapman and Hall, New York.
- Fox, J.; Weisberg, S. (2011). *An R Comparision to Applied Regression*. Sage, Thousand Oaks, CA, second edition.
- Green, P. E.; Krieger, A. M (1987). A simple Heuristic for Selecting „good“ Products in Conjoint Analysis. *Application of Management Science*, 5, 131-153.
- Green, P. E.; Rao, V. R. (1971). Conjoint Measurement for Quantifying Judgmental Data. *Journal of Marketing Research*, 8, 355-363.
- Green, P. E.; Savitz, J. (1994). Applying Conjoint Analysis to Product Assortment and Pricing in Retailing Research, *Pricing Strategy and Practice*, 4-19.
- Green, P. E.; Srinivasan, V. (1978). Conjoint Analysis in Consumer Research: Issues and Outlook. *Journal of Consumer Research*, 5, 103-123.
- Green, P. E.; Wind, Y. (1975). New Way to Measure Consumers' Judgments, *Harvard Business Review*, 53, 107-117.

- Gustafsson, A., Herrmann, A., e Huber, F. (2001). Conjoint analysis as an instrument of market research practice. In Conjoint measurement (pp. 5-46). Springer Berlin Heidelberg.
- Gustafsson, A.; Ekdahl, F.; Bergman, B. (1999). Conjoint analysis: A usefull tool in the design process. *Total Quality Management*, v. 10, n. 3, p. 327-343.
- Johnson, R. M., and K. A. Olberts. (1991). Using conjoint analysis in pricing studies: Is one price variable enough? American Marketing Association Advanced Research Technique Forum Conference Proceedings volume (4), 164–73.
- Kamakura, W. A. and Srivastava, R. K. (1984), Predicting Choice Shares Under Conditions of Brand Interdependence, *Journal of Marketing Research*, 21, 420-434.
- Koelemeijer, K., and H. Oppewal. (1999). Assessing the effects of assortment and ambience: A choice experimental approach. *Journal of Retailing* 75(3, Fall):319–45.
- Kotler, P.; Armstrong, G. (2008). *Princípios de marketing*. 12ª edição. Rio de Janeiro: Prentice Hall do Brasil LTDA. 624 p.
- Kuhfeld, W. F. (1997). Efficient Experimental Design Using Computerized Searches, *Sawtooth Software: Research Paper Series*.
- Kunsch, H. R. (1989). The jackknife and the bootstrap for general stationary observations. *The annals of Statistics*, 1217-1241.
- Liu, R. Y. (1988). Bootstrap procedures under some non-iid models. *The Annals of Statistics*, 16(4), 1696-1708.
- Louviere, J. J. and Woodworth, G. (1983), Design and Analysis of Simulated Consumer Choice or Allocation Experiments: An Approach Based on Aggregate Data, *Journal of Marketing Research*, 20, 350-367
- Luce, D. (1959), *Individual Choice Behavior*, John Wiley and Sons, New York.
- Luce, R. D. e J. W. Tukey (1964). Simultaneous Conjoint Measurement - A New Type of Fundamental Measurement, *Journal of Mathematical Psychology*, 1, 1-27.
- Macfie, H. J.; Bratchell, N.; Greenhoff, K.; Vallis, L. V. Designs to balance the effect of order of presentation and first-order carry-over effects in hall tests. *Journal of Sensory Studies*, v. 4, p. 129-148, 1989.

- Mahajan, V.; Green, P. E.; Goldberg, S. M. (1982). A Conjoint Model for Measuring Self and Cross-Price/Demand Relationships. *Journal of Marketing Research*, v. 19, n. 3, p. 334-342.
- McFadden, D. (1974). Conditional Logit Analysis of Qualitative Choice Behavior, *Frontiers in Econometrics*, Academic Press, P. Zarembka Eds., New York.
- Moore, W. L., J. J. Louviere, R. Verma. (1999). Using conjoint analysis to help design product platforms. *Journal of Product Innovation Management* 16(1, January): 27–39.
- Moskowitz, H. R.; Itty, B.; Katz, R.; Maier, A.; Beckley, J.; Flores, L. (2004). Hispanic and non-hispanic responses to concepts for four foods. *Journal of Sensory Studies*, Hoboken, v. 19, n. 6, p. 459-485.
- Oppewal, H., H. J. P. Timmermans, and J. J. Louviere. (1997). Modelling the effects of shopping centre size and store variety on consumer choice behaviour. *Environment and Planning A* 29(6):1073–90.
- Rao, V. R. Applied conjoint analysis. Springer Science & Business Media, 2014.
- Rhee, H. T.; Yang, S. B. (2015). How does hotel attribute importance vary among different travelers? An exploratory case study based on a conjoint analysis. *Electronic markets*, 25(3), 211-226.
- Rossi, E., Allenby, G., McCulloch, R. Bayesian statistics and marketing, John Wiley e Sons, 2005.
- Ryan, M., and S. Farrar. (2000). Using conjoint analysis to elicit preferences for health care. *British Medical Journal* 320:1530–33.
- Sands, S.; K. Warwick (1981). What product Benefits to offer to whom: an Application of Conjoint Segmentation, *California Management Review*, 24, 69-74.
- Silva, C. H. O.; Bastos, F. S. (2010). *Introdução à Conjoint analysis*. IX Encontro Mineiro de Estatística (MGEST), Viçosa.
- Stockwell, M. S.; Rosenthal, S. L.; Sturm, L. A.; Mays, R. M.; Bair, R. M.; Zimet, G. D. (2011). The effects of vaccine characteristics on adult women's attitudes about vaccination: A conjoint analysis study. *Vaccine*, 29(27), 4507-4511.
- Train, K. E. (2009). Discrete choice methods with simulation. Cambridge university press.

Tscheulin, D. K.; B, Helmig. (1998). The optimal Design of Hospital Advertising by Means of Conjoint Measurement. *Journal of Advertising Research*, 38, 35 - 46.

Van der Waerden, P., H. Oppewal, and H. Timmermans. (1993). Adaptive choice behavior of motorists in congested shopping centre parking lots. *Transportation* 20:395–408.

Yoo, D. I.; H. Ohta (1995). Optimal Pricing and Product Planning for new Multittribute Products based on Conjoint Analysis, *International Journal of Production Economics*, 38, 245-254.

Zimmermann, T. M., Clouth, J., Elosge, M., Heurich, M., Schneider, E., Wilhelm, S., e Wolfrath, A. (2013). Patient preferences for outcomes of depression treatment in Germany: A choice-based conjoint analysis study. *Journal of affective disorders*, 148(2), 210-219.

CAPÍTULO I

INFERÊNCIA VIA *BOOTSTRAP* NA *RATINGS-BASED CONJOINT ANALYSIS*

Resumo: No presente estudo apresenta-se uma aplicação em dados reais da *Ratings-Based Conjoint Analysis* (RBCA) com inferências adicionais para as importâncias relativas de cada atributo, realizadas por meio de suas distribuições empíricas, obtidas via método *Bootstrap*. Ajustou-se o modelo da RBCA em um estudo com três atributos e doze tratamentos (rótulos de produtos alimentícios lácteos) utilizando um esquema fatorial completo. Acessando a distribuição empírica das importâncias relativas foi possível concluir que dois dos atributos avaliados possuíam importâncias relativas equivalentes, sob o ponto de vista estatístico, inferência que até então não era obtida.

Palavras-Chave: Resíduos, notas de preferência, produtos lácteos.

Abstract: In this study we present an application in real data of Ratings-Based Conjoint Analysis (RBCA) with additional inferences for the relative importance of each attribute, made with the empirical distributions, obtained by the Bootstrap method. The RBCA model was adjusted with three attributes and twelve treatments (labels of dairy products) in a full factorial design. Accessing the empirical distribution of the relative importances, it was possible to conclude that two of the evaluated attributes had equivalent amounts, from the statistical point of view, an inference that had not been obtained until then.

Keywords: Residues, preference notes, dairy products.

1. Introdução

A origem da *Ratings-Based Conjoint Analysis* (RBCA) ou Análise Conjunta Baseada em Notas foi na área de psicologia matemática, com o trabalho teórico do psicólogo Luce e do estatístico Tukey (Luce e Tukey, 1964), que desenvolveram procedimentos para mensurar o efeito conjunto de variáveis independentes sobre uma variável dependente ordenada. No entanto, os responsáveis por introduzir essa técnica ao *Marketing*, seu principal segmento de utilização, foram Green e Rao (1971), que difundiram e ilustraram suas principais aplicações em áreas potenciais desse segmento

(Green e Rao, 1971; Green e Srinivasan, 1978; Green e Krieger, 1987; Tscheulin e Helmig, 1998).

Mais especificamente, na RBCA, notas de preferência (ou variáveis dependentes) são utilizadas para representar a utilidade (benefício ou satisfação) que um indivíduo atribui a um determinado tratamento (produto, serviço, conceito, etc.). Essa utilidade é decomposta em utilidades parciais (ou coeficientes de preferência), associadas a um nível ou a uma característica do tratamento (variáveis independentes), via modelo de regressão linear múltipla (Rao, 2014). Posteriormente, os parâmetros desse modelo são estimados e suas estimativas representam o efeito de cada nível sobre a nota média de preferência dos tratamentos. Um resultado adicional e de grande interesse são as importâncias relativas. Tais quantidades definidas por funções não lineares dos parâmetros estimados informam o percentual de importância, que os indivíduos associam a cada atributo no momento da aquisição de um tratamento (Green e Srinivasan, 1978).

Além das aplicações clássicas da RBCA no *marketing*, essa metodologia tem sido fundamental para estudos desenvolvidos na área de tecnologia de alimentos, auxiliando pesquisadores a entender a preferência dos consumidores por novos produtos alimentícios e rótulos de embalagens, bem como na melhoria de produtos já existentes, na divulgação e na elaboração de estratégias para a distribuição dos mesmos. Aplicações recentes da RBCA nesse segmento são apresentadas por Caleguer, Minim e Benassi (2007), Ares e Deliza (2010), Martínez Michel, Anders e Wismer (2011), Carneiro et al. (2012), Boesch (2013), Lima Filho et al. (2015) e em Ares et al. (2016) entre outros.

Na tentativa de aprimorar a modelagem da estrutura de preferência dos indivíduos, a RBCA tem sido adaptada e outros modelos frequentistas foram propostos, tais como: modelos de regressão por indivíduos e *cluster analysis* (Green, 1977), modelos de mistura ou de classes latentes (DeSarbo et al., 1992), modelos de regressão via mínimos quadrados parciais (Jaeger et al., 2013) e os modelos de regressão beta (Resende e Nakano, 2015). Nas duas primeiras abordagens modela-se a heterogeneidade entre a preferência dos indivíduos, enquanto nas duas últimas busca-se melhorar a qualidade do ajuste em modelos agregados, que estimam um único vetor de parâmetros para representar a preferência média dos indivíduos da amostra.

No entanto, em nenhuma dessas alternativas é possível inferir sobre a importância relativa dos atributos, isto é, testar hipóteses ou realizar comparações

estatísticas sobre tais quantidades, uma vez que apenas suas estimativas pontuais são conhecidas. Essa limitação pode ser vista em Andrade et al. (2016), que avaliou a importância dos atributos: Tipo de Corte, Apresentação, Tempero e Preço, sobre a intenção de compra de consumidores de carne de cordeiro em duas situações diferentes, a primeira quando a carne era comprada em dias regulares de semana e a segunda em finais de semanas com datas comemorativas. No entanto, não foi possível testar se as diferenças encontradas foram significativas, uma vez que a distribuição amostral das importâncias relativas não é conhecida. O mesmo ocorreu com Reis et al. (2016), que pretendiam comparar a importância dos atributos: informação nutricional, tipo de nutriente e tipo de processamento sobre a intenção de compra de sucos de romã/laranja, quando um grupo de consumidores foi submetido a uma restrição de tempo durante o processo de avaliação e o outro não. Nesses casos, a ausência de tais inferências, além de tornar a RBCA menos informativa, pode conduzir a conclusões errôneas sobre a importância relativa dos atributos, uma vez que a variabilidade de seus estimadores não é avaliada.

Neste sentido, propõe-se aqui a aplicação do método *Bootstrap* (Efron, 1979) na RBCA. Além da análise tradicional, será possível explorar a distribuição empírica das importâncias relativas e inferir sobre tais quantidades, sendo esse o principal objetivo do presente trabalho, já que sob o enfoque frequentista isso ainda não é realizado. Mais especificamente, medidas de precisão (como por exemplo, o erro padrão e o viés) poderão ser obtidas e intervalos de confiança construídos para as importâncias relativas dos atributos, no intuito de avaliar a acurácia das estimativas e realizar comparações estatísticas entre tais quantidades. Para tal, os dados utilizados são provenientes de um experimento real (Lima, 2015), conduzido com 180 indivíduos e 12 tratamentos.

2. Material e métodos

2.1 Dados experimentais analisados

O estudo envolveu $N = 180$ consumidores, de ambos os sexos, com idade entre 18 e 65 anos, sendo estudantes e funcionários da Universidade Federal Rural do Rio de Janeiro e das unidades da Embrapa Agroindústria de Alimentos e Agrobiologia. Os tratamentos avaliados foram rótulos de produtos alimentícios lácteos, definidos a partir da combinação dos níveis de $R = 3$ atributos: Produto (“Iogurte *light*”, “Queijo prato” e

“Achocolatado”), Informação (“Tabela” e “Semáforo”) e Marca (“Conhecida” e “Desconhecida”). A avaliação foi conduzida por notas e o delineamento fatorial foi o do perfil completo, com $J = 12$ tratamentos, conforme apresentados na Tabela 1 e ilustrados na Figura 1. Detalhes adicionais sobre o experimento e a coleta dos dados podem ser consultados em Lima (2015).

Tabela 1 - Delineamento dos Tratamentos em estudo.

Tratamento	Produto	Marca	Informação
1	Achocolatado	Conhecida	Tabela
2	Achocolatado	Desconhecida	Tabela
3	Achocolatado	Conhecida	Semáforo
4	Achocolatado	Desconhecida	Semáforo
5	Iogurte <i>light</i>	Conhecida	Tabela
6	Iogurte <i>light</i>	Desconhecida	Tabela
7	Iogurte <i>light</i>	Conhecida	Semáforo
8	Iogurte <i>light</i>	Desconhecida	Semáforo
9	Queijo prato	Conhecida	Tabela
10	Queijo prato	Desconhecida	Tabela
11	Queijo prato	Conhecida	Semáforo
12	Queijo prato	Desconhecida	Semáforo



Figura 1 - Tratamentos. (1) Iogurtes (2) Bebidas achocolatadas, (3) Queijos pratos ilustrados. Fonte: Lima (2015).

2.2 Ratings-Based Conjoint Analysis (RBCA)

O enfoque frequentista da RBCA foi realizado ajustando-se aos dados o modelo de regressão linear múltipla, conforme apresentado em (1). Neste modelo, as interações duplas e de ordem superior são negligenciadas por apresentarem baixa explicação da variação da variável dependente, além de difícil interpretação prática (Louviere, 1988).

$$Y_{nj} = \beta_0 + \sum_{r=1}^R \sum_{s=1}^{m_r} X_{rs}^j \beta_{rs} + \varepsilon_{nj} \quad (1)$$

Ou simplesmente:

$$Y = X\beta + \varepsilon$$

Em que:

$Y = (Y_{11} \dots Y_{1J} Y_{21} \dots Y_{2J} \dots Y_{N1} \dots Y_{NJ})^T$ é o vetor de dimensão $(NJ \times 1)$ com as notas de preferência dos J tratamentos, fornecidas pelos N indivíduos;

$X = [X_1 \ X_2 \ \dots \ X_N]^T$ é a matriz de delineamento do experimento $(NJ \times p)$, com $p = 1 + \sum_{r=1}^R m_r$, sendo m_r o número de níveis presentes em cada um dos r atributos. Adicionalmente, $X_1 = X_2 = \dots = X_N$ são matrizes de dimensão $(J \times p)$, sendo cada X_{rs}^j uma variável *dummy* que especifica o código do s -ésimo nível do r -ésimo atributo, presente no j -ésimo tratamento. Portanto, para $n = 1, 2, \dots, N$ tem-se:

$$X_n = \begin{bmatrix} 1 & X_{11}^1 & \dots & X_{1m_1}^1 & X_{21}^1 & \dots & X_{2m_2}^1 & \dots & X_{R1}^1 & \dots & X_{Rm_R}^1 \\ 1 & X_{11}^2 & \dots & X_{1m_1}^2 & X_{21}^2 & \dots & X_{2m_2}^2 & \dots & X_{R1}^2 & \dots & X_{Rm_R}^2 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \ddots & \vdots & \ddots & \vdots & \ddots & \vdots \\ 1 & X_{11}^J & \dots & X_{1m_1}^J & X_{21}^J & \dots & X_{2m_2}^J & \dots & X_{R1}^J & \dots & X_{Rm_R}^J \end{bmatrix}$$

$\beta = (\beta_0 \ \beta_{11} \dots \beta_{1m_1} \ \beta_{21} \dots \beta_{2m_2} \dots \beta_{R1} \dots \beta_{Rm_R})^T$ é o vetor de parâmetros desconhecidos (ou coeficientes de preferência) de dimensão $(p \times 1)$, em que β_{rs} representa o efeito do nível s do atributo r sobre a nota média de preferência;

$\varepsilon = (\varepsilon_{11} \dots \varepsilon_{1J} \ \varepsilon_{21} \dots \varepsilon_{2J} \dots \varepsilon_{N1} \dots \varepsilon_{NJ})^T$ é o vetor de erros aleatórios e não observáveis $(NJ \times 1)$ do modelo, associados aos N indivíduos e aos J tratamentos.

Note que a codificação proposta para os elementos da matriz X implica em um problema de multicolinearidade perfeita, o que inviabiliza as soluções clássicas de estimação por Mínimos Quadrados Ordinários (MQO) ou por Máxima Verossimilhança

(MV) (Gujarati, 2009), visto que $\mathbf{X}^T \mathbf{X}$ é uma matriz singular. Neste sentido, completa-se o seu posto impondo a restrição paramétrica para que a soma dos efeitos dos níveis de cada atributo seja nula [Kuhfeld (2003)], isto é:

$$\sum_{s=1}^{m1} \beta_{1s} = \sum_{s=1}^{m2} \beta_{2s} = \dots = \sum_{s=1}^{mR} \beta_{Rs} = 0 \quad (2)$$

Inferências sobre os níveis dos atributos são conduzidas analisando-se o sinal dos coeficientes do modelo (β_{rs}). Devido à restrição (2) pode-se interpretar que, se:

$$\begin{cases} \beta_{rs} < 0 \Rightarrow \text{O nível } s \text{ do atributo } r \text{ diminui (efeito negativo) a nota de preferência} \\ \beta_{rs} > 0 \Rightarrow \text{O nível } s \text{ do atributo } r \text{ aumenta (efeito positivo) a nota de preferência} \end{cases}$$

A importância dos atributos é avaliada pela medida definida por Green e Srinivasan (1978) em (3), denominada por importância relativa [$IR_r(\%)$]:

$$\begin{aligned} I_r &= \max_s(\beta_r) - \min_s(\beta_r) \\ I_t &= \sum_{r=1}^R I_r \\ IR_r(\%) &= \frac{I_r}{I_t} 100\% \end{aligned} \quad (3)$$

Como os valores paramétricos de β_{rs} , I_r , I_t e $IR_r(\%) \forall_{r,s}$ são desconhecidos, estes são substituídos por suas estimativas, produzidas pelos respectivos estimadores $\hat{\beta}_{rs}$, $\hat{I}_r$, $\hat{I}_t$ e $\widehat{IR}_r(\%)$, após a observação da amostra aleatória. Note que no enfoque frequentista os valores de $IR_r(\%)$ são estimados pontualmente, no entanto, não são testados estatisticamente, pois sua distribuição amostral não é conhecida.

2.3 Princípios básicos do *Bootstrap* não paramétrico

Seja $\mathbf{y} = (y_1 \ y_2 \ \dots \ y_n)^T$ o conjunto de valores observados que representa uma realização da variável aleatória $\mathbf{Y} = (Y_1 \ Y_2 \ \dots \ Y_n)^T$, proveniente de uma população com função de distribuição desconhecida e indexada pelo parâmetro θ , isto é, $F(\mathbf{y}, \theta)$. Por simplicidade, adote θ como um escalar, embora sua representação possa ser vetorial. Admita o interesse em conhecer a distribuição amostral da estatística $\hat{\theta} = t(\mathbf{y})$, utilizada para estimar $\theta = t(\mathbf{Y})$. Aplicando o método *Bootstrap* são obtidos, com reposição, K novos conjuntos de dados $\mathbf{y}_k^* = (y_1^* \ y_2^* \ \dots \ y_n^*)$, com $k = 1, 2, \dots, K$. Se para cada k -ésima reamostra ou réplica *Bootstrap* a estatística $\hat{\theta} = t(\mathbf{y})$ for estimada, isto é,

$\hat{\theta}_k^* = t(\mathbf{y}_k^*)$, pode-se acessar a distribuição *Bootstrap* de $\hat{\theta}^*$ ou a distribuição empírica de $\hat{\theta}$. Essa distribuição é considerada uma estimativa da distribuição amostral de $\hat{\theta}$ (Fox e Weisberg, 2011). Logo, histogramas, medidas de posição e de dispersão podem ser utilizadas para explorá-la e intervalos de confiança e testes de hipóteses podem ser construídos/realizados para avaliar o parâmetro θ . Em geral, as principais medidas calculadas são a média (4), a variância (5), o erro padrão da média (6) e o viés (7) *Bootstrap*:

$$\hat{\theta}_B = \overline{\hat{\theta}^*} = \hat{E}(\hat{\theta}^*) = \frac{\sum_{k=1}^K \hat{\theta}_k^*}{K} \quad (4)$$

$$\widehat{Var}(\hat{\theta}^*) = \frac{\sum_{k=1}^K (\hat{\theta}_k^* - \hat{\theta}_B)^2}{K - 1} \quad (5)$$

$$\widehat{EP}(\hat{\theta}^*) = \sqrt{\widehat{Var}(\hat{\theta}^*)} \quad (6)$$

$$\hat{B}(\hat{\theta}^*) = \hat{\theta}_B - \hat{\theta} \quad (7)$$

Existem vários intervalos de confiança *Bootstrap*. O leitor pode consultar detalhes em Efron e Tibshirani (1993) e Chernick e LaBudde (2014). Neste trabalho foi empregado apenas o intervalo Percentil *Bootstrap*, que utiliza os quantis empíricos da distribuição de $\hat{\theta}^*$ para obter os respectivos limites inferior e superior. Portanto, admita que cada $\hat{\theta}_{(1)}^*, \hat{\theta}_{(2)}^*, \dots, \hat{\theta}_{(K)}^*$, tal que $\hat{\theta}_{(1)}^* < \hat{\theta}_{(2)}^* < \dots < \hat{\theta}_{(K)}^*$, represente a k -ésima estatística de ordem da distribuição de $\hat{\theta}^*$. Logo, o intervalo com $100(1 - \alpha)\%$ de confiança para o parâmetro θ é definido conforme em (8) (Efron, 1981):

$$IC_{(1-\alpha)}(\theta) = [\hat{\theta}_{(L)}^*; \hat{\theta}_{(U)}^*] \quad (8)$$

Em que $L = \left\lfloor \frac{(K+1)\alpha}{2} \right\rfloor$, $U = \left\lfloor (K+1) \left(1 - \frac{\alpha}{2}\right) \right\rfloor$ e $\lfloor x \rfloor$ representa o maior número inteiro menor ou igual à x .

2.4 Método *Residuals Bootstrap* na RBCA

Admita que $\mathbf{y} = (y_{11} \dots y_{1J} y_{21} \dots y_{2J} \dots y_{N1} \dots y_{NJ})^T$ represente o vetor composto por NJ variáveis respostas (variáveis dependentes) ou notas de preferência observadas e que $\mathbf{X}_n^j = [1 \ X_{11}^j \dots X_{1m_1}^j \ X_{21}^j \dots X_{2m_2}^j \dots X_{R1}^j \dots X_{Rm_R}^j]$ represente o vetor linha da matriz de delineamento $\mathbf{X}$, referente ao n -ésimo indivíduo e composto pela codificação dos níveis dos atributos presentes no j -ésimo tratamento.

No método *Residual Bootstrap*, a reamostragem ocorre nos resíduos do modelo (1) ajustado. Essa abordagem é também conhecida como *Fixed-X Resampling* ou *Model-Based Resampling*, pois as variáveis explicativas são consideradas fixas, como tradicionalmente ocorre na análise de regressão (Davison e Kuonen, 2002). O algoritmo utilizado para essa abordagem é definido conforme os seguintes passos:

- i. Ajuste o modelo (1) aos dados da amostra original ou mestra;
- ii. Obtenha o vetor $\hat{\boldsymbol{\beta}} = (\hat{\beta}_0 \hat{\beta}_{11} \dots \hat{\beta}_{1m_1} \hat{\beta}_{21} \dots \hat{\beta}_{2m_2} \dots \hat{\beta}_{R1} \dots \hat{\beta}_{Rm_R})^T$ por MQO ou por MV e, posteriormente, o vetor de resíduos $\hat{\boldsymbol{\varepsilon}} = \mathbf{Y} - \hat{\mathbf{Y}} = (\hat{\varepsilon}_{11} \dots \hat{\varepsilon}_{1J} \hat{\varepsilon}_{21} \dots \hat{\varepsilon}_{2J} \dots \hat{\varepsilon}_{N1} \dots \hat{\varepsilon}_{NJ})^T$;
- iii. Realize a transformação, conforme sugerido por Davison e Hinkley (1997), para obter resíduos com média nula e mesma variância:

$$r_{nj} = \frac{\hat{\varepsilon}_{nj}}{\sqrt{1 - h_i}} - \bar{r} \sim (0, \sigma^2)$$

Em que $\bar{r} = \frac{\sum_{j=1}^J \sum_{n=1}^N r_{nj}}{NJ}$ e h_i é o i -ésimo elemento diagonal da matriz $\mathbf{H}$, de dimensão $(NJ \times NJ)$, definida como $\mathbf{H} = \mathbf{X}(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T$, em que $(\mathbf{X}^T \mathbf{X})$ é a matriz com as restrições de soma nula para o efeito dos níveis de cada atributo, logo, de posto completo e com inversa única;

- iv. Obter os NJ pares de valores $(\mathbf{X}_{nj}^*, y_{nj}^*)$, que irão compor a primeira réplica *Bootstrap*, de acordo com a relação:

$$\begin{cases} \mathbf{X}_{nj}^* = \mathbf{X}_n^j \\ y_{nj}^* = \mathbf{X}_{nj}^* \hat{\boldsymbol{\beta}} + \hat{\varepsilon}_{nj}^* \end{cases} \quad \forall \begin{matrix} n=1,2,\dots,N \\ j=1,2,\dots,J \end{matrix}$$

Em que $\hat{\varepsilon}_{nj}^*$ é um valor selecionado de forma aleatória e com reposição dos resíduos do modelo ajustado no passo i. Estime os parâmetros do modelo (1) e a importância relativa de cada atributo conforme mencionado na seção 2.1.

- v. Repita o passo iv. K vezes para acessar a distribuição empírica ou *Bootstrap* das estatísticas de interesse.

2.5 Aspectos computacionais

Utilizou-se o *software* livre R (R Core Team, 2014) para condução das análises estatísticas. A estimação dos parâmetros do modelo da RBCA ocorreu por MQO, utilizando a função *lm* do pacote *stats*. As restrições paramétricas de soma nula para os efeitos dos níveis de cada atributo foram impostas via argumento *contr.sum* da função *contrasts* do pacote MASS. O *design* de tratamentos foi realizado pelas funções *expand.grid* do pacote *Conjoint* (Bak e Bartlomowicz, 2009). O procedimento *Residual Bootstrap* foi realizado com a função *boot* do pacote *boot* (Davison e Hinkley, 1997). O intervalo de confiança Percentil *Bootstrap* foi obtido pela função *boot.ci* do referido pacote, via argumento *perc*. Os testes de Shapiro e Wilk (1965), D'Agostino (1970) e Anscombe-Glynn (1983), foram empregados a 5% de probabilidade para avaliar, respectivamente, hipóteses de normalidade, assimetria e curtose das distribuições empíricas das importâncias relativas. Os códigos em linguagem R estão disponíveis no Apêndice I.

3. Resultados e discussão

A Tabela 2 apresenta as estimativas de média, viés e erro padrão *Bootstrap* para os parâmetros do modelo da RBCA. Um total de $K = 1999$ réplicas foram utilizadas para construir a distribuição empírica de cada estimador.

Tabela 2 - Estimativas e estatísticas² fornecidas pelo método *Residuals Bootstrap*.

Atributos	Níveis	$\hat{\beta}_B$	$\hat{B}(\hat{\beta}^*)$	$\widehat{EP}(\hat{\beta}^*)$
	Intercepto	5,3454	-0,0005	0,0214 (0,0418)
Marca	Conhecida	0,0812	-0,0002	0,0215 (0,0418)
	Desconhecida	-0,0812	0,0002	0,0215 (0,0418)
Informação	Tabela	-0,0051	0,00002	0,0216 (0,0418)
	Semáforo	0,0051	-0,00002	0,0216 (0,0418)
Produto	Achocolatado	-0,8686	0,0004	0,0310 (0,0591)
	Iogurte <i>Light</i>	1,2064	-0,0007	0,0295 (0,0591)
	Queijo Prato	-0,3378	0,0003	0,0302 (0,0591)

* Valores entre parênteses representam as estimativas clássicas de erro padrão.

Os resultados práticos indicam que os indivíduos avaliados preferem produtos alimentícios lácteos cujo rótulo seja de uma marca Conhecida ($\hat{\beta}_B = 0,0812$), com

² Em que $\hat{\beta}_B$, $\hat{B}(\hat{\beta}^*)$ e $\widehat{EP}(\hat{\beta}^*)$ representam as estimativas de média, viés e erro padrão *Bootstrap*, respectivamente.

informação nutricional do tipo Semáforo ($\hat{\beta}_B = 0,0051$) e associados a um Iogurte *Light* ($\hat{\beta}_B = 1,2064$). Adicionalmente, a estimativa do intercepto sugere que a nota média de preferência dos doze tratamentos em estudo é de 5,3454.

Brown (1993) e Leão e Melo (2009) corroboram esse resultado afirmando que produtos com marcas conhecidas tendem a influenciar a percepção do consumidor, uma vez que refletem qualidade, procedência e confiança. Já Sonnemberg et al. (2013) avaliaram a influência do semáforo nutricional no rótulo de alimentos comercializados no refeitório de um hospital e verificaram que o seu efeito foi positivo, isto é, os consumidores passaram a consumir com maior frequência alimentos que apresentavam a luz verde do que alimentos com a luz vermelha. Ares, Gimenez e Gámbaro (2008) concluíram que devido à classificação de “produto saudável”, os iogurtes *light* retornaram as maiores intenções de compra.

Note ainda que as estimativas de média *Bootstrap* apresentaram baixo viés e baixo erro padrão, o que é desejável, pois demonstra que essa medida pode ser empregada como um estimador para os parâmetros de interesse. Lima (2015) encontrou resultados similares em termos de estimativas pontuais, no entanto, ajustando o modelo da RBCA, aos mesmos dados, via MQO. Vale também ressaltar que as estimativas de erro padrão *Bootstrap* foram inferiores às estimativas clássicas, o que vai de encontro com os resultados obtidos por Algama e Rasheed (2010), que verificaram em um estudo de simulação, que os valores de média *Bootstrap* coincidem com as estimativas de MQO ou de *Jackknife*, no entanto, fornecem menor viés e menor erro padrão, principalmente em pequenas amostras. Logo, se não for possível avaliar a preferência de muitos indivíduos, o que pode ocorrer na prática por questões de tempo, logística ou por indisposição dos indivíduos, à alternativa *Bootstrap* pode ser mais indicada que a abordagem frequentista tradicional, pois fornecerá estimativas mais precisas.

A Tabela 3 apresenta os intervalos Percentis *Bootstrap* para os parâmetros do modelo da RBCA. Nesta Tabela, $\hat{\beta}_{(50)}^*$ e $\hat{\beta}_{(1950)}^*$ representam as estatísticas de ordem da distribuição de cada $\hat{\beta}^*$, que acumulam 2,5% e 97,5% de probabilidade, respectivamente. Adicionalmente, admitiu-se a normalidade dos resíduos para a construção do tradicional intervalo com 95% de confiança baseado na distribuição de probabilidade t-Student³. Como critério de comparação, a amplitude de cada intervalo (*A*) também foi definida.

³ Em que *L* = limite inferior e *U* = limite superior do respectivo intervalo de confiança.

Tabela 3 - Intervalos baseados na distribuição t-Student e intervalos Percentis *Bootstrap*.

Atributos	Níveis	t-Student			Percentil <i>Bootstrap</i>		
		<i>L</i>	<i>U</i>	<i>A</i>	$\hat{\beta}_{(50)}^*$	$\hat{\beta}_{(1950)}^*$	<i>A</i>
-	Intercepto	5,264	5,428	0,164	5,302	5,388	0,087
Marca	Conhecida	-0,0005	0,1634	0,164	0,039	0,125	0,086
	Desconhecida	-0,1634	0,0005	0,164	-0,125	-0,039	0,086
Informação	Tabela	-0,087	0,077	0,164	-0,048	0,037	0,085
	Semáforo	-0,077	0,087	0,164	-0,037	0,048	0,085
Produto	Achocolatado	-0,984	-0,752	0,232	-0,930	-0,805	0,125
	Iogurte <i>light</i>	1,090	1,322	0,232	1,149	1,263	0,113
	Queijo prato	-0,453	-0,222	0,231	-0,397	-0,280	0,117

Pode-se verificar que o intervalo Percentil *Bootstrap* obteve amplitude relativamente inferior à do intervalo baseado na distribuição t-Student, isto é, o método proposto forneceu intervalos mais curtos e com o mesmo nível de confiança, o que é desejável. Uma possível explicação para tal fato é que na RBCA as notas de preferência não possuem exatamente distribuição normal, pois são discretas, como no presente estudo, ou contínuas em uma escala entre 1 e 9, o que pode ter afetado as inferências clássicas (Resende e Nakano, 2015). Como no intervalo Percentil *Bootstrap* as estatísticas de ordem das distribuições empíricas são utilizadas para obter os limites de confiança (Efron, 1981), essa possível ausência ou o enfraquecimento da suposição de normalidade não compromete as inferências.⁴

Ao avaliar os intervalos Percentis *Bootstrap* pode-se inferir sobre a significância e a diferença estatística entre o efeito dos níveis do atributo Marca e do atributo Produto, e também sobre a nulidade e a igualdade estatística do efeito dos níveis do atributo Informação. Portanto, mesmo que o coeficiente estimado para o nível Semáforo tenha sido positivo (Tabela 2), seu efeito é nulo e sua presença não aumenta a nota média de preferência de um tratamento. Tal conclusão é obtida observando se o valor zero pertence ao respectivo intervalo. É interessante notar que se a inferência fosse conduzida via intervalo clássico, a conclusão sobre os níveis do atributo Marca seria diferente, uma vez que o valor zero agora pertence a esses intervalos mais longos e que se sobrepõe. A Figura 2 ilustra os histogramas de probabilidade dessas distribuições. As linhas em vermelho indicam os limites inferior e superior dos intervalos Percentis *Bootstrap* e as linhas em azul os limites do intervalo baseado na distribuição t-Student.

⁴ Adicionalmente, tais intervalos não exigem a definição de quantidades pivotais, o que é uma vantagem, visto que essa tarefa nem sempre é trivial.

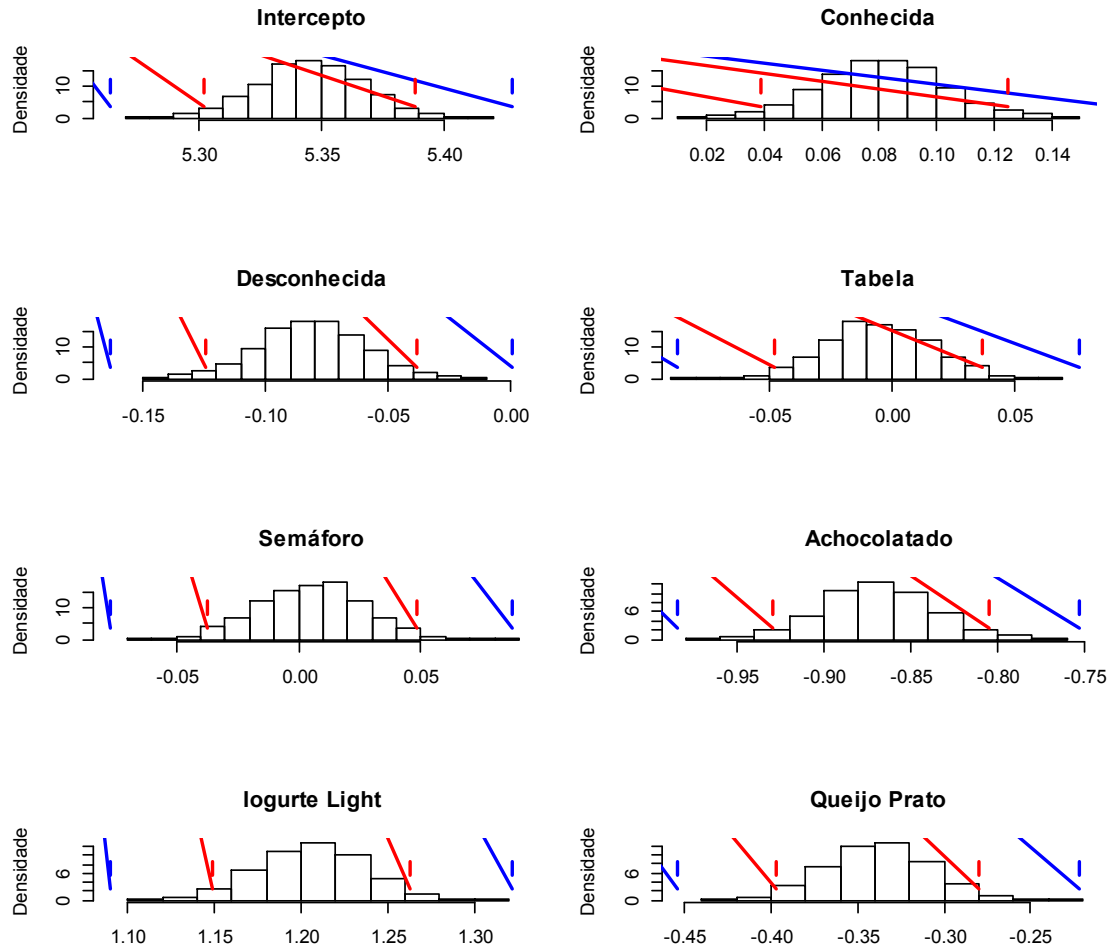


Figura 2 - Histograma das distribuições empíricas dos estimadores do modelo da RBCA.

A Tabela 4 apresenta as estimativas de importância relativa para os atributos em estudo e as demais estatísticas fornecidas pelo método *Residual Bootstrap*.

Tabela 4 - Estimativas e estatísticas *Bootstrap* para as importâncias relativas dos atributos.

Atributos	$\widehat{IR}_B(\%)$	$\widehat{B}(\widehat{IR}^*)$	$\widehat{EP}(\widehat{IR}^*)$	$\widehat{IR}_{(50)}^*$	$\widehat{IR}_{(1950)}^*$
Marca	7,1228	-0,1280	1,7709	3,525	10,619
Informação	1,5550	1,0992	1,1494	0,062	4,212
Produto	91,3221	-0,9712	2,0465	87,231	95,288

O atributo Produto obteve a maior importância relativa [$\widehat{IR}_B(\%) = 91,3221$], segundo a amostra avaliada. Posteriormente, o atributo Marca [$\widehat{IR}_B(\%) = 7,1228$] e o atributo Informação nutricional [$\widehat{IR}_B(\%) = 1,5550$]. Conforme Lima (2015), tais resultados demonstram que os indivíduos parecem já ter avaliado os alimentos com um conhecimento prévio sobre os mesmos. Por isso o atributo Marca e o atributo Informação nutricional obtiveram uma importância relativa tão baixa, em relação ao atributo Produto. A autora também encontrou resultados similares ao proceder a RBCA

da maneira tradicional. Curiosamente, essa mesma discrepância foi verificada quando a análise foi segmentada em dois grupos ou *clusters*.

Ao analisar os intervalos Percentis *Bootstrap* conclui-se que a importância relativa do atributo Produto é estatisticamente diferente da importância relativa dos atributos Informação e Marca, uma vez que os intervalos não se sobrepõem. No entanto, isso já era esperado, devido a grande diferença entre os valores pontuais estimados. O resultado mais importante refere-se à igualdade estatística entre a importância relativa dos atributos Marca e Informação. Essa conclusão não era até então obtida na análise frequentista tradicional da RBCA, pois os intervalos de confiança não eram conhecidos. Logo, no contexto empresarial, essas inferências podem oferecer maior confiabilidade aos gestores no momento de tomar uma decisão ou de fazer um investimento, pois agora tais conclusões não serão baseadas apenas em diferenças matemáticas e sim em informações probabilísticas. A Figura 3 apresenta o Histograma de probabilidade com as das distribuições empíricas das importâncias relativas.

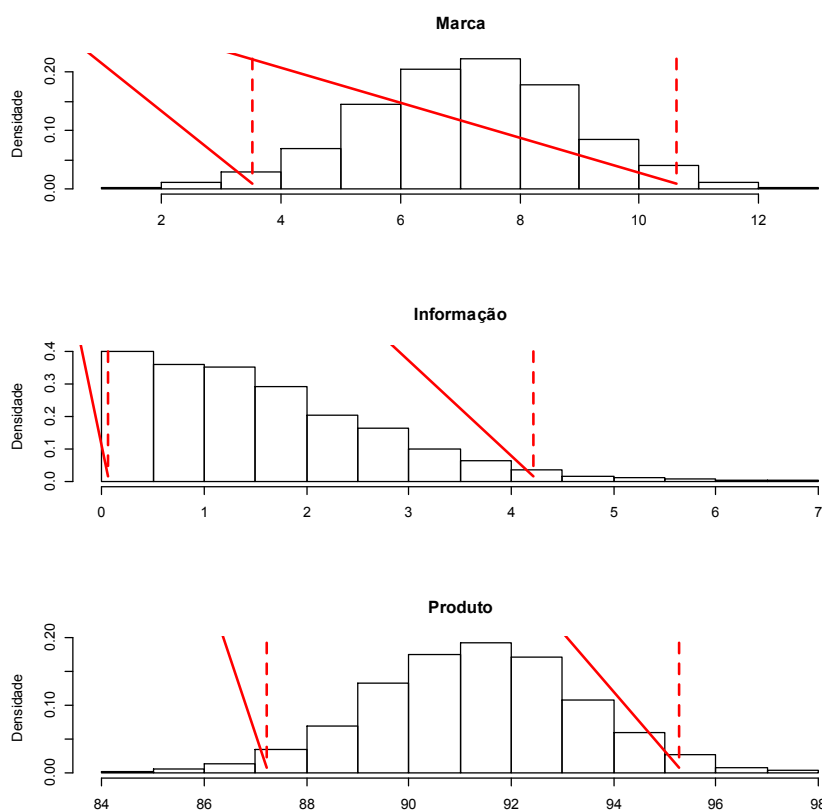


Figura 3 - Histograma das distribuições empíricas das importâncias relativas.

Destaca-se que com exceção do atributo Informação Nutricional (valor-p = 0,000), o teste de Shapiro-Wilk indicou que a distribuição empírica das importâncias relativas é gaussiana. Note que a assimetria presente na distribuição da importância

relativa do atributo Informação nutricional foi a principal causa da ausência de normalidade, conforme indica o teste de D'Agostino (valor-p = 0,000), embora o teste de Anscombe-Glynn tenha também indicado o excesso de curtose (valor-p = 0,000).

4. Conclusões

Além de oferecer estimativas pontuais mais precisas, o método *Residual Bootstrap* permitiu avaliar a distribuição empírica das importâncias relativas dos atributos e inferir estatisticamente sobre essas quantidades, o que sob o enfoque frequentista até então não era realizado, devido à dificuldade de se derivar medidas de erro padrão ou de acessar a distribuição amostral dos respectivos estimadores. Logo, foi possível concluir que o atributo Marca e Informação Nutricional possuem, sob o ponto de vista estatístico, importâncias relativas equivalentes. Tal conclusão poderia auxiliar gestores no momento de tomar uma decisão comercial, isto é, em quais atributos investir ou modificar no *design* do produto.

5. Referências

- Algamal, Z. Y.; e Rasheed, K. B. (2010). Re-sampling in Linear Regression Model Using Jackknife and Bootstrap. *Iraqi Journal of Statistical Science* 18, 59-73.
- Andrade, J. C., Nalério, É. S., Giongo, C., de Barcellos, M. D., Ares, G., e Deliza, R. (2016). Influence of evoked contexts on rating-based conjoint analysis: Case study with lamb meat. *Food Quality and Preference*, 53, 168-175.
- Anscombe, F.J.; Glynn, W.J. (1983) Distribution of kurtosis statistic for normal statistics. *Biometrika*, 70, 1, 227-234.
- Ares, G., e Deliza, R. (2010). Studying the influence of package shape and colour on consumer expectations of milk desserts using word association and conjoint analysis. *Food Quality and Preference*, 21(8), 930-937.
- Ares, G., Arrúa, A., Antúnez, L., Vidal, L., Machín, L., Martínez, J., ... e Giménez, A. (2016). Influence of label design on children's perception of two snack foods: Comparison of rating and choice-based conjoint analysis. *Food Quality and Preference*, 53, 1-8.

- Ares, G.; Gimenez, A.; Gámbaro, A (2008). Understanding consumers' perception of conventional and functional yogurts using word association and hard laddering. *Food Quality and Preference*, 19, 635-643.
- Bak, A.; Bartlomowicz, T. (2009). Conjoint analysis method and its implementation in conjoint R package. *Wroclaw: University of Economics*.
- Boesch, I. (2013). Preferences of processing companies for attributes of Swiss milk: A conjoint analysis in a business-to-business market. *Journal of dairy science*, 96(4), 2183-2189.
- Brown, S. (1993). Postmodern marketing? *European Journal of Marketing*, 27, v. 4, 19-34.
- Caleguer, V. F., Minim, V. P. R., e Benassi, M. T. (2007). Impact of package on the consumer purchase intention for a powdered orange flavoured soft drink. *Braz. J. Food Technol*, 10, 159-168.
- Carneiro, J. D. D. S., Minim, V. P. R., Silva, C. H. O., Regazzi, A. J., Chaves, J. B. P., e Minim, L. A. (2012). Sugarcane spirit market share simulation: an application of conjoint analysis. *Food Science and Technology (Campinas)*, 32(4), 645-652.
- Chernick, M. R.; LaBudde, R. A. (2014). *An introduction to bootstrap methods with applications to R*. John Wiley e Sons.
- D'Agostino, R.B. (1970). Transformation to Normality of the Null Distribution of G1. *Biometrika*, 57, 3, 679-681.
- Davison, A. C.; Hinkley, D. V. (1997). *Bootstrap Methods and their Applications*. Cambridge University Press, Cambridge.
- Davison, A. C.; Kuonen, D. (2002). An Introduction to the Bootstrap with Applications in R. *Statistical Computing and Statistical Graphics Newsletter*, 13(1), 6-11.
- DeSarbo, W. S.; Wedel, M.; Vriens, M.; Ramaswamy, V. (1992), Latent Class Metric Conjoint Analysis, *Marketing Letters*, 3, 273-289.
- Efron, B. (1979). Bootstrap methods: another look at the jackknife. *Annals of Statistics*, 7:1 – 26.
- Efron, B. (1981). Nonparametric standard errors and condence intervals (with discussion). *Canadian Journal of Statistic*, 9, 139-172.

- Efron, B.; Tibshirani, R. J. (1993). *An Introduction to the Bootstrap*. Chapman and Hall, New York.
- Fox, J.; Weisberg, S. (2011). *An R Comparision to Applied Regression*. Sage, Thousand Oaks, CA, second edition.
- Green, P. E. (1977). A New Approach to Market Segmentation. *Business Horizons*, 20, 61-73.
- Green, P. E.; Krieger, A. M (1987). A simple Heuristic for Selecting ‘good’ Products in Conjoint Analysis. *Application of Management Science*, 5, 131-153.
- Green, P. E.; Rao, V. R. (1971). Conjoint Measurement for Quantifying Judgmental Data. *Journal of Marketing Research*, 8, 355-363.
- Green, P. E.; Srinivasan, V. (1978). Conjoint Analysis in Consumer Research: Issues and Outlook. *Journal of Consumer Research*, 5, 103-123.
- Gujarati, D. N. (2009). *Basic econometrics*. Tata McGraw-Hill Education.
- Jaeger, S. R.; Mielby, L. H.; Heymann, H.; Jia; Frøst, M. B. (2013). Analysing conjoint data with OLS and PLS regression: a case study with wine. *Journal of the Science of Food and Agriculture*, 93(15), 3682-3690.
- Kuhfeld, W. F. (2003). *Conjoint Analysis. SAS Institute Publications*. Disponível em: < <http://citeseerx.ist.psu.edu/viewdoc/download?doi=10.1.1.194.8152erep=rep1etype=pdf> >.
- Leão, A. L. M. S.; Mello, S. C. B. (2009). “Valor de marca” para quem? Rumo a uma teoria da significação das marcas pelos consumidores. *ROC*, n.10, v.5.
- Lima Filho, T., Della Lucia, S. M., Lima, R. M., e Minim, V. P. R. (2015). Conjoint analysis as a tool to identify improvements in the packaging for irradiated strawberries. *Food Research International*, 72, 126-132.
- Lima, M. F. (2015). Semáforo Nutricional (Traffic Light Labeling) e seu efeito na avaliação dos alimentos pelo consumidor. (Dissertação), Mestrado em Ciência e Tecnologia de Alimentos, Instituto de Tecnologia, Universidade Federal Rural do Rio de Janeiro, 2015.
- Louviere, J. J. (1988). Conjoint analysis modelling of stated preferences: a review of theory, methods, recent developments and external validity. *Journal of transport economics and policy*, 93-119.

- Luce, R. D.; Tukey, J. W. (1964). Simultaneous Conjoint Measurement - A New Type of Fundamental Measurement. *Journal of Mathematical Psychology*, 1, 1-27.
- Martínez Michel, L., Anders, S., e Wismer, W. V. (2011). Consumer Preferences and willingness to pay for value- added chicken product attributes. *Journal of food science*, 76(8).
- R Core Team (2017). R: A language and environment for statistical computing. *R Foundation for Statistical Computing*, Vienna, Austria (<http://www.r-project.org>).
- Rao, V. R. (2014). *Applied conjoint analysis*. New York, NY: Springer.
- Reis, F.; Machín, L.; Rosenthal, A.; Deliza, R.; Ares, G. (2016). Does a time constraint modify results from rating-based conjoint analysis? Case study with orange/pomegranate juice bottles. *Food Research International*, 90, 244-250.
- Resende, V. S.; Nakano, E. Y. (2015). Análise Conjunta Baseada em Notas via Modelo de Regressão Beta. *Revista Brasileira de Biometria*, 33(1), 51-62.
- Shapiro, S. S.; Wilk, M. B. (1965). An analysis of variance test for normality (complete samples). *Biometrika*, 591-611.
- Sonnenberg, L.; Gelsomin, E.; Levy, D. E.; Riis, J.; Barraclough, S.; Thorndike, A. N. (2013). A traffic light food labeling intervention increases consumer awareness of health and healthy choices at the point-of-purchase. *Preventive medicine*, 57(4), 253-257.
- Tscheulin, D. K.; B, Helmig. (1998). The optimal Design of Hospital Advertising by Means of Conjoint Measurement. *Journal of Advertising Research*, 38, 35 - 46.

Apêndice I

```
#### Pacotes

# install.packages("boot")
# install.packages("conjoint")
install.packages("moments")

library("boot")
library("conjoint")
library("moments")

#### Banco de dados

# Criando o Design

experiment <- expand.grid(
  Marca <- c("Conhecida", "Desconhecida"),
  Semaforo <- c("Tabela", "Semaforo"),
  Produto <- c("Achocolatado", "Iogurte", "Queijo"))
design <- caFactorialDesign(data = experiment, type = "full")
design

# Replicando o design para "n" consumidores

n <- 179
p <- 12

design_final <- do.call("rbind", replicate(n, design, simplify = F))
design_final
colnames(design_final) <- c("Marca", "Semaforo", "Produto")
head(design_final)

# Notas dos Tratamentos

setwd("c:/DadosR")
Notas <- stack(read.table("notas_tese_A2.txt", h = T, dec = ","))[,1]
head(Notas)

# Impondo o contraste de Soma zero

contrasts(design_final$Marca) <- contr.sum
contrasts(design_final$Semaforo) <- contr.sum
contrasts(design_final$Produto) <- contr.sum

# Banco de dados final

dados <- data.frame(design_final, Notas)
head(dados)

# Análise Tradicional por MQO

reg <- lm(Notas ~ Marca + Semaforo + Produto, data = dados)
summary(reg)
head(model.matrix(reg), p)

# Verificando o efeito de qual nível foi estimado

contrasts(design_final$Marca)
contrasts(design_final$Semaforo)
contrasts(design_final$Produto)
```

```

# Coeficientes Estimados para cada nível

(b0 <- reg$coef[1])
(b_Marca <- contrasts(design_final$Marca)%*%reg$coef[2]) # Marca
(b_Semaforo <- contrasts(design_final$Semaforo)%*%reg$coef[3]) # Semaforo
(b_Produto <- contrasts(design_final$Produto)%*%reg$coef[c(4,5)]) # Produto
(betas <- rbind(b0, b_Marca, b_Semaforo, b_Produto))

# Erro Padrão tradicional

ep_reg <- cbind(coef(summary(reg))[ , 2])
(ep <- cbind(c(ep_reg[1], rep(ep_reg[2],4), rep(ep_reg[4],3))))

#### Intervalos de Confiança t-Student

(L <- round(betas - qt(0.975, (n*p - 8 - 1), lower.tail = T)*ep, 7))
(U <- round(betas + qt(0.975, (n*p - 8 - 1), lower.tail = T)*ep, 7))
(IC <- cbind(L,U))

# Inferências (Importâncias e Importâncias Relativas)

(I_Marca <- max(b_Marca)- min(b_Marca))
(I_Semaforo <- max(b_Semaforo)- min(b_Semaforo))
(I_Produto <- max(b_Produto)- min(b_Produto))

(I_Total <- I_Marca + I_Semaforo + I_Produto)
(I <- rbind(I_Marca, I_Semaforo, I_Produto))

(IR_Marca <- (I_Marca/I_Total)*100)
(IR_Semaforo <- (I_Semaforo/I_Total)*100)
(IR_Produto <- (I_Produto/I_Total)*100)

(IR_Total <- IR_Marca + IR_Semaforo + IR_Produto)
(IR <- rbind(IR_Marca, IR_Semaforo, IR_Produto))

##### Residuals Bootstrap - Fixed-X resampling

reg <- lm(Notas ~ Marca + Semaforo + Produto, data = dados)
X = model.matrix(reg) # Matriz de delineamenro do modelo de Regressão
fitr = fitted(reg) # Valores Ajustados da Regressão
reg_glm <- glm(Notas ~ Marca + Semaforo + Produto, family = "gaussian", data = dados)
r = glm.diag(reg_glm)$rp - mean(glm.diag(reg_glm)$rp) # resíduos com média nula e variância constante

f <- function(formula, data, indices)
{

y_star <- fitr + r[indices]
fit <- lm(y_star ~ X - 1)

b0 <- fit$coef[1]
b_Marca <- contrasts(design_final$Marca)%*%fit$coef[2] # Marca
b_Semaforo <- contrasts(design_final$Semaforo)%*%fit$coef[3] # Semaforo
b_Produto <- contrasts(design_final$Produto)%*%fit$coef[c(4,5)] # Produto
betas <- c(b0, b_Marca, b_Semaforo, b_Produto)

I_Marca <- max(b_Marca)- min(b_Marca) # Importância Marca
I_Semaforo <- max(b_Semaforo)- min(b_Semaforo) # Importância Semaforo
I_Produto <- max(b_Produto)- min(b_Produto) # Importância Produto
I_Total <- I_Marca + I_Semaforo + I_Produto

IR_Marca <- (I_Marca/I_Total)*100 # Importância Relativa Marca

```

```

IR_Semaforo <- (I_Semaforo/I_Total)*100 # Importância Relativa Semaforo
IR_Produto <- (I_Produto/I_Total)*100 # Importância Relativa Produto
IR <- c(IR_Marca, IR_Semaforo, IR_Produto)

return(c(betas, IR))
}

result <- boot(data = dados, statistic = f, R = 1999, formula = Notas ~ Marca + Semaforo + Produto)
result

### Médias Bootstrap

(media_boot <- cbind(colMeans(result$t)))

### Intervalos de Confiança Percentil Bootstrap

k = 11
percentil <- matrix(44,11,4)
colnames(percentil) <- c("Q_L", "Q_S", "LI", "LS")

for (i in 1:k)
{
percentil[i,] <- c(boot.ci(result, conf = 0.95, index = i, typ = "perc")$percent[, -1])
}
percentil

### Gráficos da distribuição empírica dos betas

par(mfrow=c(4,2))

hist(result$t[,1], prob = T, main = "Intercepto", xlab = "", ylab = "Densidade", xlim = c(IC[1,1], IC[1,2]))
abline(v = c(IC[1,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[1,c(3,4)]), col = "red", lty = 2, lwd = 2, xlim = c(IC[2,1], IC[2,2]))
hist(result$t[,2], prob = T, main = "Conhecida", xlab = "", ylab = "Densidade")
abline(v = c(IC[2,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[2,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,3], prob = T, main = "Desconhecida", xlab = "", ylab = "Densidade", xlim = c(IC[3,1], IC[3,2]))
abline(v = c(IC[3,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[3,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,4], prob = T, main = "Tabela", xlab = "", ylab = "Densidade", xlim = c(IC[4,1], IC[4,2]))
abline(v = c(IC[4,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[4,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,5], prob = T, main = "Semáforo", xlab = "", ylab = "Densidade", xlim = c(IC[5,1], IC[5,2]))
abline(v = c(IC[5,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[5,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,6], prob = T, main = "Achocolatado", xlab = "", ylab = "Densidade", xlim = c(IC[6,1], IC[6,2]))
abline(v = c(IC[6,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[6,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,7], prob = T, main = "Iogurte Light", xlab = "", ylab = "Densidade", xlim = c(IC[7,1], IC[7,2]))
abline(v = c(IC[7,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[7,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,8], prob = T, main = "Queijo Prato", xlab = "", ylab = "Densidade", xlim = c(IC[8,1], IC[8,2]))
abline(v = c(IC[8,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[8,c(3,4)]), col = "red", lty = 2, lwd = 2)

# Gráficos da distribuição empírica das Importâncias Relativas

par(mfrow=c(3,1))

```

```

hist(result$T[,9], prob = T, main = "Marca", xlab = "Marca", ylab = "Densidade")
abline(v = c(percentil[9,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$T[,10], prob = T, main = "Informação", xlab = "Informação", ylab = "Densidade")
abline(v = c(percentil[10,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$T[,11], prob = T, main = "Produto", xlab = "Produto", ylab = "Densidade")
abline(v = c(percentil[11,c(3,4)]), col = "red", lty = 2, lwd = 2)

```

Teste de Normalidade, Assimetria e Curtose para os parâmetros e Importâncias relativas

```

k = 11
teste <- matrix(0,11,8)
colnames(teste) <- c("W", "p", "Ass", "z", "p", "Curt.", "z", "p")
rownames(teste) <- c("b0", "b11", "b12", "b21", "b22", "b31", "b32", "b33", "IR1", "IR2", "IR3")

for (i in 1:k)
{
teste[i,1] <- round(shapiro.test(result$T[,i])$statistic,4)
teste[i,2] <- round(shapiro.test(result$T[,i])$p.value,4)
teste[i,3] <- round(agostino.test(result$T[,i])$statistic[1],4)
teste[i,4] <- round(agostino.test(result$T[,i])$statistic[2],4)
teste[i,5] <- round(agostino.test(result$T[,i])$p.value,4)
teste[i,6] <- round(anscombe.test(result$T[,i])$statistic[1],4)
teste[i,7] <- round(anscombe.test(result$T[,i])$statistic[2],4)
teste[i,8] <- round(anscombe.test(result$T[,i])$p.value,4)
}
teste

```

CAPÍTULO II

INFERÊNCIA VIA BOOTSTRAP NA CHOICE-BASED CONJOINT ANALYSIS

Resumo: No presente trabalho apresenta-se uma aplicação da *Choice-Based Conjoint Analysis* (CBCA) com inferências adicionais para as Probabilidades de Escolha dos tratamentos e para a Razão de Escolha de cada atributo, realizadas por meio de suas distribuições empíricas, obtidas via método *Bootstrap*. O estudo envolveu três atributos e oito tratamentos (iogurtes *light* sabor morango) em um esquema fatorial completo. Acessando a distribuição empírica de tais quantidades foi possível inferir sobre as mesmas, o que no contexto frequentista não era até então realizado. Mais especificamente, medidas de viés e de erro padrão foram obtidas para as probabilidades e razões de escolha, o que permitiu avaliar a acurácia dessas estimativas. Adicionalmente, intervalos de confiança foram construídos e tais quantidades puderam ser estatisticamente comparadas.

Palavras-Chave: *Bootstrap* por Pares, iogurte *light*, software livre R.

Abstract: In this paper we present an application of Choice-Based Conjoint Analysis (CBCA) with additional inferences for the Choice of Treatment Probabilities and for Choice Ratio of each attribute, made through its empirical distributions, obtained through Bootstrap. The study involved three attributes and eight treatments (light strawberry flavor yoghurts) in a complete factorial scheme. Accessing an empirical distribution of such volumes was possible to infer as part of it, what not classical context was not until realized. More specifically, bias and standard error measures were obtained for probabilities and ratios of choice, which allow us to evaluate the firm's estimates. In addition, confidence intervals were constructed and such volumes could be statistically compared.

Keywords: Paired Bootstrap, light yoghurt, free software R.

1. Introdução

A origem da *Choice-Based Conjoint Analysis* (CBCA) ou Análise Conjunta Baseada em Escolhas foi com o trabalho de Louviere e Woodworth (1983), que combinaram os conceitos experimentais da *Conjoint Analysis* tradicional (Luce e Tukey, 1964) com os modelos econométricos de escolha discreta, mais especificamente,

o modelo Logit Multinomial, inicialmente derivado por Luce (1959) e completamente especificado por McFadden (1974). Posteriormente, na área de *marketing*, seu principal segmento de utilização, outros trabalhos foram desenvolvidos [Kamakura e Srivastava (1984), Johnson e Olberts (1991), Koelemeijer e Oppewal (1999), Moore, Louviere e Verma (1999)], o que difundiu e consolidou essa metodologia.

Mais especificamente, na CBCA, as escolhas dos indivíduos (ou variáveis dependentes) são utilizadas para representar a utilidade (benefício ou satisfação) atribuída a um determinado tratamento (produto, serviço, conceito, etc.). Essa utilidade é decomposta em utilidades parciais (ou coeficientes de preferência), associadas a um atributo ou a uma característica do tratamento (variáveis independentes), via modelo de regressão logit (McFadden, 1974). Posteriormente, os parâmetros desse modelo são estimados e suas estimativas utilizadas para calcular as probabilidades de escolha do conjunto de tratamentos. Uma vez que o modelo empregado é não linear, a interpretação dos parâmetros estimados não é trivial e, portanto, um resultado de grande importância é a Razão de Escolha dos atributos, a qual informa o quanto o nível de um determinado atributo é mais provável de ser escolhido em relação aos demais (Greene, 2003).

Além das aplicações clássicas no *marketing*, a CBCA tem sido empregada em outras áreas do conhecimento, em especial na tecnologia de alimentos. Essa metodologia tem sido utilizada para avaliar a preferência de consumidores por novos produtos alimentícios e rótulos de embalagens, uma vez que essa alternativa é mais realista e prática no que se refere a situações reais de compra, pois o consumidor não precisa atribuir uma nota para cada um dos tratamentos, o que é bem menos desgastante (Moore, 2004). O leitor pode consultar algumas aplicações nessa área em Lockshin et al. (2006), Della Lucia et al. (2010), Deliza et al. (2010), Tempesta et al. (2010), Szücs et al. (2014) e Meyerding (2015).

No entanto, muitos pesquisadores ainda adotam a abordagem da *Conjoint Analysis* Baseada em Notas (RBCA) para avaliar a preferência do consumidor, especialmente no Brasil, alegando que essa alternativa é mais informativa, uma vez que um conjunto de notas fornece mais informação do que uma única escolha [Moore, Gray-Lee e Louviere (1998), Karniouchina et al. (2009) e Asioli et al. (2016)]. Neste sentido, outras abordagens frequentistas foram propostas para tentar popularizar a CBCA e aprimorar a modelagem da estrutura de preferência, como o modelo Dogit Multinomial (Gaudry e Dagenais, 1977), que permite acomodar grupos de indivíduos cativos a uma mesma decisão (indivíduos cuja preferência independe das alternativas de tratamentos) e o modelo Logit Sequencial (McFadden, 1978), que possibilita que as decisões ou

escolhas sejam realizadas de forma hierárquica. Adicionalmente, os modelos de classes latentes (DeSarbo, Ramaswamy e Cohen, 1995) e modelo Logit com Heterogeneidade Ajustada (Chesher e Santos Silva, 2002) permitem avaliar a influência da heterogeneidade entre a escolha dos indivíduos.

Contudo, em nenhuma dessas alternativas é possível inferir sobre a probabilidade dos tratamentos e sobre a razão de escolha dos atributos, isto é, testar hipóteses ou realizar comparações estatísticas a respeito de tais quantidades, uma vez que apenas suas estimativas pontuais são conhecidas. Logo, a ausência de tais inferências, além de tornar a CBCA menos informativa, pode conduzir a conclusões errôneas, já que a variabilidade dos respectivos estimadores não é avaliada. Mais especificamente, tratamentos com probabilidades estatisticamente iguais poderão ser admitidos como probabilisticamente diferentes, uma vez que apenas uma estimativa pontual é informada. Naturalmente, se o custo associado a cada um destes tratamentos for relativamente diferente, uma decisão errônea poderia ocasionar custos desnecessários ao gestor.

Neste sentido, propõe-se aqui a aplicação do método *Bootstrap* (Efron, 1979) na CBCA. Além da análise tradicional, será possível explorar a distribuição empírica das probabilidades/razões de escolha e inferir sobre tais quantidades, sendo esse o principal objetivo do presente trabalho, já que sob o enfoque frequentista isso ainda não é realizado. Para tal, os dados utilizados são provenientes de um experimento real (Della Lucia et al., 2010), conduzido com 144 indivíduos e 8 tratamentos.

2. Material e Métodos

2.1 Dados Experimentais Analisados

O presente estudo envolveu $N = 144$ indivíduos que residiam na cidade de Viçosa (Minas Gerais, Brasil) e $R = 3$ atributos: informação sobre o conteúdo de Açúcar (“açúcar”), informação sobre o conteúdo de Gordura (“gordura”) e informação sobre o conteúdo de Proteína (“proteína”), com os respectivos níveis: açúcar (“0% de açúcar” e “com adoçante”), gordura (“0% de gordura” e “baixo teor de gordura”) e proteína (“enriquecido com proteínas do soro do leite” e “enriquecido com proteínas bioativas”). O delineamento fatorial empregado foi o do perfil completo, com $J = 8$

tratamentos, conforme apresentado na Tabela 1. Maiores detalhes sobre a condução do experimento podem ser consultados em Della Lucia et al. (2010).

Tabela 1 – Delineamento dos tratamentos em estudo.

Tratamento	Açúcar	Gordura	Proteína
1	0% de açúcar	0% de gordura	Com proteínas do soro do leite
2	Com adoçante	0% de gordura	Com proteínas do soro do leite
3	0% de açúcar	Baixo teor de gordura	Com proteínas do soro do leite
4	Com adoçante	Baixo teor de gordura	Com proteínas do soro do leite
5	0% de açúcar	0% de gordura	Com proteínas bioativas
6	Com adoçante	0% de gordura	Com proteínas bioativas
7	0% de açúcar	Baixo teor de gordura	Com proteínas bioativas
8	Com adoçante	Baixo teor de gordura	Com proteínas bioativas

2.2 Choice-Based Conjoint Analysis (CBCA)

Na CBCA, a relação ou o modelo para utilidade aleatória é a definida em (1).

$$U = X\beta + \varepsilon \quad (1)$$

Em que:

$U = [U_{11} \dots U_{1J} U_{21} \dots U_{2J} \dots U_{N1} \dots U_{NJ}]^T$ é o vetor $(NJ \times 1)$ de utilidades aleatórias associadas aos J tratamentos, atribuídas pelos N indivíduos;

$X = [X_1 \ X_2 \ \dots \ X_N]^T$ em que $X_1 = X_2 = \dots = X_N$ são matrizes de dimensão $(J \times R)$ que especificam a codificação do s -ésimo nível do r -ésimo atributo, presente no j -ésimo tratamento. Para $n = 1, 2, \dots, N$ tem-se que:

$$X_n = \begin{bmatrix} X_1^1 & X_2^1 & \dots & X_R^1 \\ X_1^2 & X_2^2 & \dots & X_R^2 \\ \vdots & \vdots & \dots & \vdots \\ X_1^J & X_2^J & \dots & X_R^J \end{bmatrix}$$

$\beta = (\beta_1 \ \beta_2 \ \dots \ \beta_R)^T$ é o vetor de parâmetros desconhecidos (ou coeficientes de preferência) de dimensão $(R \times 1)$, em que β_r representa o efeito do r -ésimo atributo sobre a utilidade aleatória média.

$\varepsilon = (\varepsilon_{11} \dots \varepsilon_{1J} \ \varepsilon_{21} \dots \varepsilon_{2J} \dots \varepsilon_{N1} \dots \varepsilon_{NJ})^T$ é o vetor $(NJ \times 1)$ de erros aleatórios e não observáveis do modelo, cuja suposição é de que cada ε_{nj} seja independente e possua

distribuição de valores extremos do tipo I ou Gumbel, com média nula e variância constante de $\frac{\pi^2}{6}$, valor este definido para que o modelo seja identificável⁵ (Train, 2009).

O modelo (2) ou Logit Multinomial (McFadden, 1974) é derivado supondo que o n -ésimo indivíduo escolherá o j -ésimo tratamento, se e somente se, $U_{nj} > U_{nk}$, $\forall j \neq k; j, k \in \{1, 2, \dots, J\}$. Portanto:

$$P(Y_n = j|\mathbf{X}) = \frac{e^{X_j\beta}}{\sum_{k=1}^J e^{X_k\beta}} \quad (2)$$

Em que $P(Y_n = j|\mathbf{X})$ é a probabilidade do j -ésimo tratamento ser escolhido pelo n -ésimo indivíduo. Como a matriz $\mathbf{X}$ é a mesma para $n = 1, 2, \dots, N$, tem-se que $P(Y_n = j|\mathbf{X}) = P(Y = j|\mathbf{X})$. Os estimadores dos parâmetros do modelo (2) são obtidos por máxima verossimilhança, maximizando o logaritmo natural de $L(\beta)$ em (3), isto é, $l(\beta) = \ln[L(\beta)]$, via método numérico de Newton-Raphson (Gallant, 2009).

$$l(\beta) = \ln \left[\prod_{j=1}^J \left(\frac{e^{X_j\beta}}{\sum_{k=1}^J e^{X_k\beta}} \right)^{n_j} \right] \quad (3)$$

Em que n_j é a variável que indica o número de vezes que o j -ésimo tratamento foi escolhido. A contribuição ou o efeito dos níveis de cada atributo sobre a escolha dos indivíduos é avaliada pela Razão de Escolha (Greene, 2003), definida conforme em (4).

$$RE_r(\mathbf{X}_p, \mathbf{X}_q) = \frac{P(Y = p|\mathbf{X})}{P(Y = q|\mathbf{X})} = e^{(x_r^p - x_r^q)\beta_r}, \quad \forall r=1, \dots, R \quad (4)$$

Em que o tratamento q é obtido do tratamento p , fixando-se $(r - 1)$ níveis dos r atributos, isto é, alterando apenas o nível de uma única característica. Logo, pode-se interpretar que, se:

$$\begin{cases} RE > 1 \Rightarrow P(Y = p|\mathbf{X}) > P(Y = q|\mathbf{X}) \\ RE = 1 \Rightarrow P(Y = p|\mathbf{X}) = P(Y = q|\mathbf{X}) \\ RE < 1 \Rightarrow P(Y = p|\mathbf{X}) < P(Y = q|\mathbf{X}) \end{cases}$$

Note que $C = \binom{m}{2}$ Razões de Escolha serão estimadas para um atributo com $m \geq 2$ níveis, em que $(.)$ indica uma combinação simples. Destaca-se que a RE é constante e depende apenas das características dos tratamentos p e q , por isso a remoção ou inclusão de outros tratamentos não altera sua proporcionalidade. Essa é uma

⁵ Um modelo é globalmente identificável se para um único conjunto de estimativas existe um valor máximo para a função de verossimilhança. Caso contrário, restrições devem ser impostas a alguns parâmetros para que o modelo se torne identificável.

característica particular do modelo Logit Multinomial, denominada por I.I.A. (*Independence of Irrelevant Alternatives*) e assegurada pela independência dos erros aleatórios (Train, 2009).

2.3 Princípios Básicos do *Bootstrap* não paramétrico

Seja $\mathbf{y} = (y_1 \ y_2 \ \dots \ y_n)^T$ o conjunto de valores observados que representa uma realização da variável aleatória $\mathbf{Y} = (Y_1 \ Y_2 \ \dots \ Y_n)^T$, proveniente de uma população com função de distribuição desconhecida e indexada pelo parâmetro θ , isto é, $F(\mathbf{y}, \theta)$. Por simplicidade, adote θ como um escalar, embora sua representação possa ser vetorial. Admita o interesse em conhecer a distribuição amostral da estatística $\hat{\theta} = t(\mathbf{y})$, utilizada para estimar $\theta = t(\mathbf{Y})$. Aplicando o método *Bootstrap* são obtidos, com reposição, K novos conjuntos de dados $\mathbf{y}_k^* = (y_1^* \ y_2^* \ \dots \ y_n^*)$, com $k = 1, 2, \dots, K$. Se para cada k -ésima reamostra ou réplica *Bootstrap* a estatística $\hat{\theta} = t(\mathbf{y})$ for estimada, isto é, $\hat{\theta}_k^* = t(\mathbf{y}_k^*)$, pode-se acessar a distribuição *Bootstrap* de $\hat{\theta}^*$ ou a distribuição empírica de $\hat{\theta}$. Essa distribuição é considerada uma estimativa da distribuição amostral de $\hat{\theta}$ (Fox e Weisberg, 2011). Logo, histogramas, medidas de posição e de dispersão podem ser utilizadas para explorá-la e intervalos de confiança e testes de hipóteses podem ser construídos/realizados para avaliar o parâmetro θ . Em geral, as principais medidas calculadas são a média (5), a variância (6), o erro padrão da média (7) e o viés (8) *Bootstrap*:

$$\hat{\theta}_B = \overline{\hat{\theta}^*} = \hat{E}(\hat{\theta}^*) = \frac{\sum_{k=1}^K \hat{\theta}_k^*}{K} \quad (5)$$

$$\widehat{Var}(\hat{\theta}^*) = \frac{\sum_{k=1}^K (\hat{\theta}_k^* - \hat{\theta}_B)^2}{K - 1} \quad (6)$$

$$\widehat{EP}(\hat{\theta}^*) = \sqrt{\widehat{Var}(\hat{\theta}^*)} \quad (7)$$

$$\hat{B}(\hat{\theta}^*) = \hat{\theta}_B - \hat{\theta} \quad (8)$$

Existem vários intervalos de confiança *Bootstrap*. O leitor pode consultar detalhes em Efron e Tibshirani (1993) e Chernick e LaBudde (2014). Neste trabalho foi empregado apenas o intervalo Percentil *Bootstrap*, que utiliza os quantis empíricos da distribuição de $\hat{\theta}^*$ para obter os respectivos limites inferior e superior. Portanto, admita

que cada $\hat{\theta}_{(1)}^*, \hat{\theta}_{(2)}^*, \dots, \hat{\theta}_{(K)}^*$, tal que $\hat{\theta}_{(1)}^* < \hat{\theta}_{(2)}^* < \dots < \hat{\theta}_{(K)}^*$, represente a k -ésima estatística de ordem da distribuição de $\hat{\theta}^*$. Logo, o intervalo com $100(1 - \alpha)\%$ de confiança para o parâmetro θ é definido conforme em (9) (Efron, 1981):

$$IC_{(1-\alpha)}(\theta) = [\hat{\theta}_{(L)}^*; \hat{\theta}_{(U)}^*] \quad (9)$$

Em que $L = \left\lfloor \frac{(K+1)\alpha}{2} \right\rfloor$, $U = \left\lfloor (K+1) \left(1 - \frac{\alpha}{2}\right) \right\rfloor$ e $\lfloor x \rfloor$ representa o maior número inteiro menor ou igual à x .

2.4 Método *Paired Bootstrap* na CBCA

Admita que $\mathbf{y} = (y_{11} \dots y_{1J} y_{21} \dots y_{2J} \dots y_{N1} \dots y_{NJ})^T$ represente o vetor composto por NJ variáveis respostas (variáveis dependentes) ou escolhas observadas e que $\mathbf{X}_n^j = [X_1^j \ X_2^j \ \dots \ X_R^j]$, represente o vetor linha da matriz de delineamento $\mathbf{X}$, referente ao n -ésimo indivíduo e composto pela codificação dos níveis presentes no j -ésimo tratamento. No método *Paired Bootstrap* a reamostragem ocorre em pares $(\mathbf{X}_{nj}^*; y_{nj}^*)$ e admite-se que cada um desses pares provém de uma distribuição conjunta ou bivariada de probabilidades $(\mathbf{X}; \mathbf{Y})$. Por isso tal abordagem é também conhecida como *Random-X Resampling* ou *Cases Resampling* (Davison e Hinkley's, 1997). O algoritmo utilizado é definido conforme os seguintes passos (Efron e Tibshirani, 1993):

- i. Selecione aleatoriamente e com reposição os NJ pares de valores $(\mathbf{X}_{nj}^*; y_{nj}^*)$, para $n = 1, 2, \dots, N$ e $j = 1, 2, \dots, J$, que irão compor a primeira réplica *Bootstrap*, de acordo com a seguinte relação;

$$\begin{cases} \mathbf{X}_{nj}^* = \mathbf{X}_n^j \\ y_{nj}^* = y_{nj} \end{cases}$$

- ii. Ajuste o modelo Logit Multinomial ao conjunto de dados obtidos no passo i., estime seus parâmetros, as Probabilidades dos tratamentos e a Razão de Escolha de cada atributo conforme mencionado na seção 2.2;
- iii. Repita os passos i. e ii. K vezes para obter a distribuição empírica ou *Bootstrap* das estatísticas de interesse.

2.5 Aspectos Computacionais

Utilizou-se o *software* livre R (R Core Team, 2017) para condução das análises estatísticas. O procedimento de estimação de parâmetros do modelo Logit Multinomial ocorreu pela função *clogit* do pacote *survival* (Therneau, 2012). O *design* de tratamentos foi realizado pelas funções *expand.grid* e *caEncodedDesign* do pacote *Conjoint* (Bak e Bartlomowicz, 2009). O procedimento *Paired Bootstrap* foi realizado pela função *boot* do pacote *boot* (Davison e Hinkley, 1997). O intervalo de confiança Percentil *Bootstrap* foi obtido pela função *boot.ci* do mesmo pacote, via argumento *perc*. O teste de Shapiro e Wilk (1965) foi empregado a 5% de probabilidade para avaliar hipóteses de normalidade sobre as distribuições empíricas das Probabilidades dos tratamentos e das Razões de Escolha dos atributos. Os códigos em linguagem R estão disponíveis no Apêndice I.

3. Resultados e Discussão

A Tabela 2 apresenta as estimativas de média, viés e erro padrão *Bootstrap* para os parâmetros do modelo da CBCA. Um total de $K = 1999$ réplicas foram utilizadas para construir a distribuição empírica de cada estimador.

Tabela 2 – Estimativas e estatísticas⁶ fornecidas pelo método *Paired Bootstrap*.

Atributos	$\hat{\beta}_B$	$\hat{B}(\hat{\beta}^*)$	$\hat{EP}(\hat{\beta}^*)$
Açúcar	-1,6759	-0,0156	0,2181 (0,2274)
Proteína	1,5241	0,0115	0,2061 (0,2166)
Gordura	-0,7918	-0,0033	0,1602 (0,1797)

* Valores entre parênteses representam as estimativas clássicas de erro padrão.

Note que as estimativas de média *Bootstrap* apresentaram baixo viés e baixo erro padrão, o que é desejável, pois demonstra que essa medida pode ser empregada como um estimador para o efeito dos atributos em estudo. Della Lúcia et al. (2010) encontraram resultados similares ao proceder a CBCA nos mesmos dados, no entanto, estimando os parâmetros do modelo Logit Multinomial por Máxima Verossimilhança (MV). O método *Bootstrap* forneceu estimativas de erro padrão muito próximas das tradicionais, no entanto, Sonh e Menke (2002) afirmam que especificamente para modelos não lineares essa a abordagem proposta é a mais indicada para quantificar a

⁶ Em que $\hat{\beta}_B$, $\hat{B}(\hat{\beta}^*)$ e $\hat{EP}(\hat{\beta}^*)$ representam as estimativas de média, viés e erro padrão *Bootstrap*, respectivamente.

incerteza das estimativas, uma vez que as aproximações lineares, tradicionalmente empregadas, podem ser inapropriadas na presença de *outliers*, dispersão significativa dos dados ou pequenas amostras.

A Tabela 3 apresenta os intervalos Percentis *Bootstrap* para os parâmetros do modelo da CBCA. Nesta Tabela, $\hat{\beta}_{(50)}^*$ e $\hat{\beta}_{(1950)}^*$ representam as estatísticas de ordem da distribuição empírica de cada $\hat{\beta}^*$, que acumulam 2,5% e 97,5% de probabilidade, respectivamente. Adicionalmente, o tradicional intervalo com 95% de confiança baseado na distribuição Normal Padrão⁷ é apresentado. Como critério de comparação, a amplitude destes intervalos (*A*) foi definida.

Tabela 3 – Intervalos baseados na distribuição Normal Padrão e intervalos Percentis *Bootstrap*.

Atributos	Percentil <i>Bootstrap</i>			Normal Padrão		
	$\hat{\beta}_{(50)}^*$	$\hat{\beta}_{(1950)}^*$	<i>A</i>	<i>L</i>	<i>U</i>	<i>A</i>
Açúcar	-2,1377	-1,2833	0,8544	-2,1061	-1,2144	0,8917
Proteína	1,1514	1,9836	0,8322	1,0879	1,9372	0,8492
Gordura	-1,1108	-0,4875	0,6233	-1,1408	-0,4360	0,7047

Com os intervalos de confiança Percentis *Bootstrap* pode-se inferir sobre a significância estatística dos atributos Açúcar, Proteína e Gordura. Tal conclusão é obtida avaliando se o valor zero pertence aos respectivos intervalos. Esse resultado coincide com o encontrado por Della Lúcia et al. (2010), no entanto, via teste de Wald (1949), cuja hipótese estatística de nulidade ($H_0: \beta = 0$) foi rejeitada a 5% de probabilidade para os três atributos em estudo. Portanto, conclui-se que as informações referentes ao conteúdo de açúcar, proteína e gordura, presentes na embalagem do iogurte *light* sabor morango, influenciam a escolha dos indivíduos.

Pode-se também verificar que o intervalo Percentil *Bootstrap* forneceu amplitude muito próxima à do intervalo clássico. Esse é um resultado esperado, uma vez que a distribuição amostral do estimador de Máxima Verossimilhança converge assintoticamente para uma distribuição normal, isto é, $\sqrt{Np}(\hat{\beta} - \beta) \xrightarrow{d} N[0, (-H)^{-1}]$ em que *H* e $(-H)$ representam a matriz Hessiana e a matriz Informação de Fisher esperada, respectivamente (Train, 2009). Logo, é intuitivo que as distribuições empíricas também se comportem de forma similar. No entanto, o teste de Shapiro-Wilk rejeitou a hipótese de normalidade para ambas as distribuições *Bootstrap* (valor-p ≈ 0). A Figura 1 apresenta os respectivos histogramas de probabilidade. As linhas em

⁷ Em que *L* = limite inferior e *U* = limite superior do respectivo intervalo.

vermelho indicam os limites inferior e superior dos intervalos Percentis *Bootstrap* e as linhas em azul os limites do intervalo baseado na distribuição Normal Padrão.

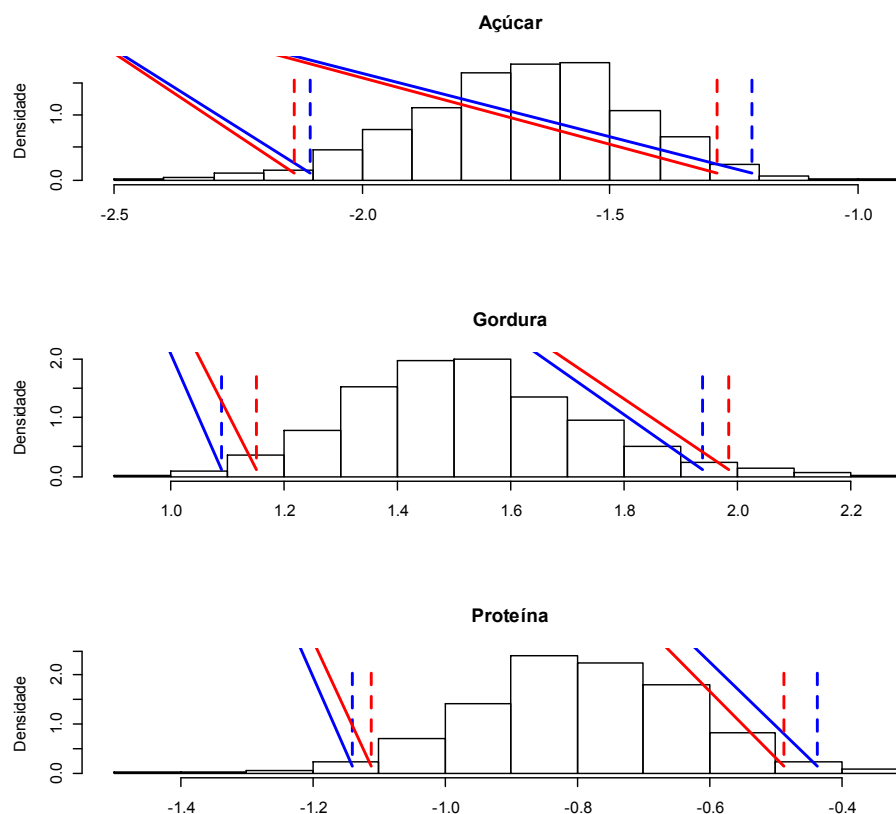


Figura 1 – Histograma das distribuições empíricas dos estimadores do modelo da CBCA.

A Tabela 4 apresenta as estimativas para as Probabilidades de Escolha do conjunto de tratamentos avaliados, além das demais estatísticas fornecidas pelo método *Paired Bootstrap*.

Tabela 4 - Estimativas e estatísticas *Bootstrap* para as Probabilidades de Escolha.

Tratamento	$\hat{P}_B$	$\hat{B}(\hat{P}_j^*)$	$\hat{EP}(\hat{P}_j^*)$	$\hat{P}_{j(50)}^*$	$\hat{P}_{j(1950)}^*$
1	0,1043	0,00001	0,0177	0,0687	0,1400
2	0,0199	0,00006	0,0052	0,0109	0,0316
3	0,0475	0,00012	0,0096	0,0295	0,0679
4	0,0091	0,00007	0,0026	0,0047	0,0149
5	0,4731	-0,00024	0,0354	0,4056	0,5421
6	0,0899	-0,00014	0,0165	0,0590	0,1230
7	0,2152	0,00006	0,0253	0,1681	0,2676
8	0,0409	0,00004	0,0087	0,0248	0,0593
	1				

Conclui-se que a maior Probabilidade de Escolha (0,4731) está associada ao tratamento 5, com características: 0% de açúcar, 0% de gordura e proteínas bioativas. Já o tratamento 4, com adoçante, baixo teor de gordura e proteínas do soro de leite, foi o que apresentou a menor Probabilidade de Escolha (0,0091). Della Lucia (2008) verificou o mesmo resultado via *Conjoint Analysis* Baseada em Notas, estimando coeficientes de preferência positivos (que aumentavam a nota de preferência) para os níveis presentes no tratamento 5 e coeficientes de preferência negativos (que diminuía a nota de preferência) para os níveis presentes no tratamento 4. Curiosamente, os mesmos resultados foram verificados quando a análise foi segmentada em dois grupos ou *clusters* de consumidores.

Além do baixo viés, as estimativas de média *Bootstrap* para as probabilidades de escolha também atendem os axiomas probabilísticos de Kolmogorov (1950), o que é de fundamental importância. Adicionalmente, com os intervalos Percentis *Bootstrap* pode-se concluir, por exemplo, que a Probabilidade de Escolha do tratamento 5 e do tratamento 7 são estatisticamente diferentes, já que não há sobreposição dos limites desses intervalos. Por outro lado, verifica-se a igualdade estatística entre a Probabilidade de Escolha do tratamento 1 e do tratamento 6. Note ainda que outras 26 comparações similares poderiam ser realizadas duas a duas. Logo, no contexto empresarial, essas inferências adicionais podem oferecer maior confiabilidade aos gestores no momento de tomar uma decisão ou de fazer um investimento, já que agora as conclusões não serão baseadas apenas em diferenças matemáticas e sim em informações probabilísticas. A Figura 2 apresenta o histograma de probabilidade com as respectivas distribuições empíricas. Segundo o teste de Shapiro-Wilk, somente a distribuição da Probabilidade de Escolha do tratamento 7 é considerada normal (valor-p = 0,051).

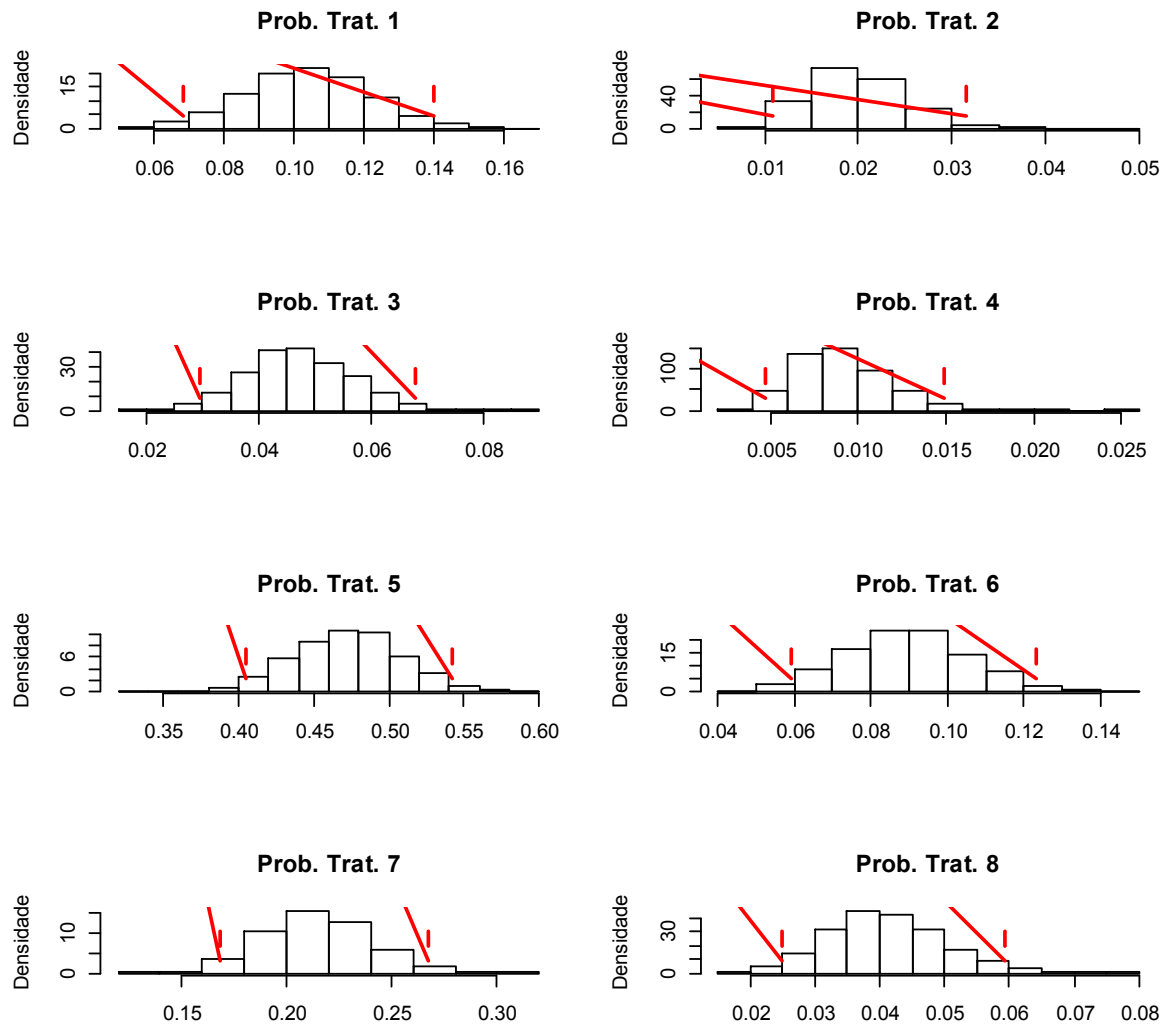


Figura 2 – Histograma das distribuições empíricas das Probabilidades de Escolha.

Para inferir sobre o efeito dos níveis de cada um dos atributos sobre a escolha do consumidor, a Tabela 5 apresenta as estimativas de Razão de Escolha e os resultados fornecidos pelo método *Paired Bootstrap*.

Tabela 5 - Estimativas e estatísticas *Bootstrap* para a Razão de Escolhas dos atributos.

Atributos	$\widehat{RE}_B$	$\widehat{B}(\widehat{RE}^*)$	$\widehat{EP}(\widehat{RE}^*)$	$\widehat{RE}_{(50)}^*$	$\widehat{RE}_{(1950)}^*$
Açúcar	5,5217	0,2608	1,3858	3,5609	8,8897
Gordura	2,2631	0,0031	0,3749	1,6515	3,1034
Proteína	4,7090	0,1705	1,0252	3,1670	7,2269

Para o atributo Açúcar, a estimativa de média *Bootstrap* da Razão de Escolha demonstra que um tratamento com “0% de Açúcar” é 5,5217 vezes mais provável de ser escolhido do que um tratamento “com adoçante”. Já para o atributo Gordura, um

tratamento com “0% de gordura” é 2,2631 vezes mais provável de ser escolhido do que um tratamento com “baixo teor de gordura” e para o atributo Proteína, um tratamento com “proteínas bioativa” é 4,7090 vezes mais provável de ser escolhido do um tratamento com “proteínas do soro do leite”.

Destaca-se que Reis (2007) verificou, via *Conjoint Analysis* Baseada em Notas, que o nível “Com adoçante” também causou impacto negativo sob a nota de preferência dos consumidores em relação ao nível “Sem adição de açúcar”. Por outro lado, os consumidores entrevistados parecem desejar consumir produtos isentos de gordura. Della Lucia (2010) explicou que o aspecto negativo associado ao termo “proteínas do soro do leite” pode estar relacionado com as características sensoriais pouco agradáveis desse nível, diferente do termo “proteínas bioativas”, associado a um alimento mais saudável, o que justifica sua maior aceitação.

A análise dos intervalos Percentis *Bootstrap* para a Razão de Escolha corrobora a significância estatística dos atributos, já verificada anteriormente, pois o valor unitário não se encontra dentro dos respectivos limites de cada intervalo. Em outras palavras, é como se testássemos $H_0: RE_r = 1 \forall_r$ contra $H_1: RE_r \neq 1 \forall_r$, decidindo então sobre a rejeição da hipótese nula para $r = 1, 2$ e 3 . Adicionalmente, nos estudos em que pelo menos um dos atributos for composto por mais de dois níveis, será possível avaliar as razões de escolha dentro de cada atributo⁸, já que testes estatísticos, como o de Wald (1949), tradicionalmente aplicados na análise frequentista da CBCA, informam apenas se o efeito do atributo é ou não diferente de zero.

A Figura 3 ilustra o histograma de probabilidade das distribuições empíricas dessas quantidades. O teste de Shapiro-Wilk rejeitou a hipótese de normalidade para ambas as Razões de Escolha (valor- $p \approx 0$), visto que tais distribuições são claramente assimétricas à direita.

⁸ Um atributo com m níveis fornece $\binom{m}{2}$ razões de escolha, em que $(.)$ representa uma combinação simples. Logo, se $m = 2$ tem-se uma única RE que compara os níveis $\{a_1, a_2\}$ de um atributo A . Já se $m = 3$ então o número de RE é 3, pois os níveis do atributo A serão comparados 2 a 2, isto é, $\{(a_1, a_2), (a_1, a_3), (a_2, a_3)\}$.

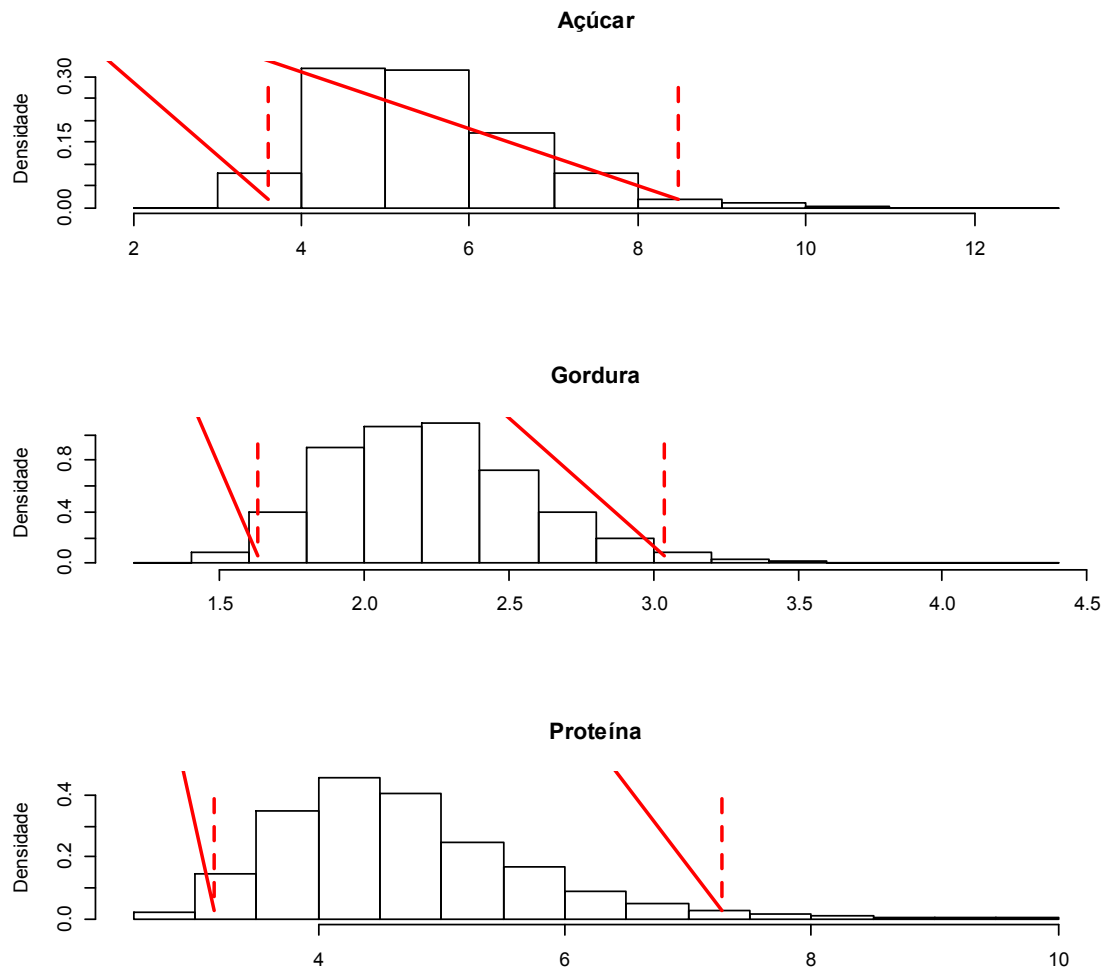


Figura 3 – Histograma das distribuições empíricas da Razão de Escolha dos atributos.

4. Conclusões

Além de oferecer estimativas pontuais mais precisas, o método *Paired Bootstrap* permitiu avaliar a distribuição empírica das probabilidades dos tratamentos e das razões de escolha e inferir estatisticamente sobre essas quantidades, o que sob o enfoque frequentista até então não era realizado, devido à dificuldade de se derivar medidas de erro padrão ou acessar a distribuição amostral dos respectivos estimadores. Logo, foi possível concluir sobre a superioridade e equivalência estatística entre a probabilidade de escolha de alguns tratamentos. Tal conclusão poderia auxiliar gestores no momento de tomar uma decisão comercial, isto é, em quais tratamentos investir ou retirar do mercado.

5. Referências Bibliográficas

- Adamowicz, W. P., M. Boxall, and M. Williams. (1998). Stated preference approaches for measuring passive use values: Choice experiments versus contingent valuation. *American Journal of Agricultural Economics* 80:64–75.
- Anscombe, F.J., Glynn, W.J. (1983) Distribution of kurtosis statistic for normal statistics. *Biometrika*, 70, 1, 227-234.
- Ares, G., Arrúa, A., Antúnez, L., Vidal, L., Machín, L., Martínez, J., Giménez, A. (2016). Influence of label design on children's perception of two snack foods: Comparison of rating and choice-based conjoint analysis. *Food Quality and Preference*, 53, 1-8.
- Asioli, D., Naes, T., Øvrum, A., Almli, V. L. (2016). Comparison of rating-based and choice-based conjoint analysis models. A case study based on preferences for iced coffee in Norway. *Food Quality and Preference*, 48, 174-184.
- Bak, A., e Bartlomowicz, T. (2009). Conjoint analysis method and its implementation in conjoint R package. *Wroclaw: University of Economics*.
- Brandley, M. A., and H. Gunn. (1990). Stated preference analysis of values of travel time in the Netherlands. *Transportation Research Record* 1285:78–89.
- Bullock, C. H., D. A. Elston, and N. A. Chalmers. (1998). An application of economic experiments to a land-use, deer hunting, and landscape change in the Scottish highlands. *Journal of Environmental Management* 52:335–51.
- Casella, G.; R.L. Berger. *Statistical Inference*. Pacific Grove, CA: Duxbury, 2002.
- Chernick, M. R., e LaBudde, R. A. (2014). *An introduction to bootstrap methods with applications to R*. John Wiley e Sons.
- Chesher, A., e Santos Silva, J. M. (2002). Taste variation in discrete choice models. *The Review of Economic Studies*, 69(1), 147-168.
- D'Agostino, R.B. (1970). Transformation to Normality of the Null Distribution of G1. *Biometrika*, 57, 3, 679-681.
- Davison, A. C., e Kuonen, D. (2002). An Introduction to the Bootstrap with Applications in R. *Statistical Computing and Statistical Graphics Newsletter*, 13(1), 6-11.

- Davison, A. C.; Hinkley, D. V. (1997). *Bootstrap Methods and their Applications*. Cambridge University Press, Cambridge.
- Deliza, R., Rosenthal, A., Hedderley, D., e Jaeger, S. R. (2010). Consumer perception of irradiated fruit: a case study using choice- based conjoint analysis. *Journal of Sensory Studies*, 25(2), 184-200.
- Della Lucia, S. M., Minim, V. P. R., Silva, C. H. O., Minim, L. A., e Silva, R. C. S. N. (2010). Análise conjunta de fatores baseada em escolhas no estudo da embalagem de iogurte light sabor morango. *Brazilian Journal of Food technology*, 11-18.
- Della Lucia, S. M.; (2008). Métodos estatísticos para a avaliação da influência de características não sensoriais na aceitação, intenção de compra e escolha do consumidor. (Tese), Doutorado em Ciência e Tecnologia de Alimentos, Departamento de Tecnologia de Alimentos, Universidade Federal de Viçosa, 2008.
- DeSarbo, W. S., Ramaswamy, V., e Cohen, S. H. (1995). Market segmentation with choice-based conjoint analysis. *Marketing Letters*, 6(2), 137-147.
- Dorfman, R. (1938), “A Note on the δ -Method for Finding Variance Formulae,” *The Biometric Bulletin*, 1, 129–137.
- Efron, B. (1979). Bootstrap methods: another look at the jackknife. *Annals of Statistics*, 7:1 – 26.
- Efron, B. (1981). Nonparametric standard errors and condence intervals (with discussion). *Canad. J. Statist.* 9, 139-172.
- Efron, B.; Tibshirani, R. J. (1993). *Na Introduction to the Bootstrap*. Chapman and Hall, New York.
- Fox, J.; Weisberg, S. (2011). *Na R Comparision to Applied Regression*. Sage, Thousand Oaks, CA, second edition.
- Gallant, A. R. (2009). *Nonlinear statistical models* (Vol. 310). John Wiley e Sons.
- Gaudry, M. J. I. e Dagenais, M. G. (1977), “The Dogit model”, *Transportation Research, Part B*, Vol. 13B, pp. 105-111.
- Greene, W. H. (2003). *Econometric analysis*. Pearson Education India.
- Haider, W., and G. O. Ewing. (1990). A model of tourist choices of hypothetical Caribbean destination. *Leisure Sciences* 12:33–47.

- Johnson, R. M., and K. A. Olberts. (1991). Using conjoint analysis in pricing studies: Is one price variable enough? American Marketing Association Advanced Research Technique Forum Conference Proceedings volume (4), 164–73.
- Kamakura, W. A. and Srivastava, R. K. (1984), Predicting Choice Shares Under Conditions of Brand Interdependence, *Journal of Marketing Research*, 21, 420-434.
- Karniouchina, E. V., Moore, W. L., van der Rhee, B., e Verma, R. (2009). Issues in the use of ratings-based versus choice-based conjoint analysis in operations management research. *European Journal of Operational Research*, 197(1), 340-348.
- Koelemeijer, K., and H. Oppewal. (1999). Assessing the effects of assortment and ambience: A choice experimental approach. *Journal of Retailing* 75(3, Fall):319–45.
- Kolmogorov, A. N. (1950). *Foundations of the Theory of Probability*.
- Lockshin, L., Jarvis, W., d’Hauteville, F., e Perrouy, J. P. (2006). Using simulations from discrete choice experiments to measure consumer sensitivity to brand, region, price, and awards in wine choice. *Food quality and preference*, 17(3), 166-178.
- Louviere, J. J. and Woodworth, G. (1983), Design and Analysis of Simulated Consumer Choice or Allocation Experiments: An Approach Based on Aggregate Data, *Journal of Marketing Research*, 20, 350-367.
- Luce, D. (1959), *Individual Choice Behavior*, John Wiley and Sons, New York.
- Luce, R. D. e J. W. Tukey (1964). Simultaneous Conjoint Measurement - A New Type of Fundamental Measurement, *Journal of Mathematical Psychology*, 1, 1-27.
- MacFie, H. J., Bratchell, N., Greenhoff, K., e Vallis, L. V. (1989). Designs to balance the effect of order of presentation and first- order carry- over effects in hall tests. *Journal of sensory studies*, 4(2), 129-148.
- McFadden, D (1974). Conditional Logit Analysis of Qualitative Choice Behavior, *Frontiers in Econometrics*, Academic Press, P. Zarembka Eds., New York, p. 105-142, 1974.
- McFadden, D. (1978), Modeling the choice of residential location, in A.Karlqvist, L.Lundqvist, F.Snicksars and J.Weibull, eds, ‘Spatial Interaction Theory and Planning Models’, North-Holland, Amsterdam, pp. 75–96.
- Meyerding, S. (2015, May). Determination of part-worth-utilities of food-labels using the choice-based-conjoint-analysis using the example of tomatoes in Germany. In *XVIII*

International Symposium on Horticultural Economics and Management 1132(pp. 17-24).

Moore, W. L. (2004). A cross-validity comparison of rating-based and choice-based conjoint analysis models. *International Journal of Research in Marketing*, 21(3), 299-312.

Moore, W. L., Gray-Lee, J., e Louviere, J. J. (1998). A cross-validity comparison of conjoint analysis and choice models at different levels of aggregation. *Marketing Letters*, 9(2), 195-207.

Moore, W. L., J. J. Louviere, R. Verma. (1999). Using conjoint analysis to help design product platforms. *Journal of Product Innovation Management* 16(1, January): 27–39.

Oppewal, H., H. J. P. Timmermans, and J. J. Louviere. (1997). Modelling the effects of shopping centre size and store variety on consumer choice behaviour. *Environment and Planning A* 29(6):1073–90.

R Core Team (2017). R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria (<http://www.r-project.org>).

Rao, V. R. (2014). *Applied conjoint analysis* (p. 56). New York, NY: Springer.

Reis, R. C. Iogurte Light Sabor Morango: Equivalência de Doçura, Caracterização Sensorial e Impacto da Embalagem na Intenção de Compra do Consumidor. 2007. 128 f. Tese (Doutorado em Ciência e Tecnologia de Alimentos)-Universidade Federal de Viçosa, Viçosa, 2007

Ryan, M., and S. Farrar. (2000). Using conjoint analysis to elicit preferences for health care. *British Medical Journal* 320:1530–33.

Shapiro, S. S., e Wilk, M. B. (1965). An analysis of variance test for normality (complete samples). *Biometrika*, 591-611.

Szűcs, V., Guerrero, L., Claret, A., Tarcea, M., Szabó, E., e Bánáti, D. (2014). Food additives and consumer preferences: a cross-cultural choice based conjoint analysis. *Acta Alimentaria*, 43(Supplement 1), 180-187.

Tempesta, T., Giancristofaro, R. A., Corain, L., Salmaso, L., Tomasi, D., e Boatto, V. (2010). The importance of landscape in wine quality perception: An integrated approach using choice-based conjoint analysis and combination-based permutation tests. *Food Quality and Preference*, 21(7), 827-836.

- Therneau T. (2012). survival: A Package for Survival Analysis in S. R package version 2.36- 12, URL <http://CRAN.R-project.org/package=survival>.
- Train, K. E. (2009). Discrete choice methods with simulation. Cambridge university press.
- Utz, K. S., Hoog, J., Wentrup, A., Berg, S., Lämmer, A., Jainsch, B., ..., e Schenk, T. (2014). Patient preferences for disease-modifying drugs in multiple sclerosis therapy: a choice-based conjoint analysis. *Therapeutic advances in neurological disorders*, 7(6), 263-275.
- Van der Waerden, P., H. Oppewal, and H. Timmermans. (1993). Adaptive choice behavior of motorists in congested shopping centre parking lots. *Transportation* 20:395–408.
- Varian, H. R. (2014). Intermediate Microeconomics: A Modern Approach: Ninth International Student Edition. WW Norton e Company.
- Ver Hoef, J. M. (2012). Who invented the delta method?. *The American Statistician*, 66(2), 124-127.
- Wald, A. (1943). Tests of statistical hypotheses concerning several parameters when the number of observations is large. *Trans. Amer. Math. Soc.*, 54, 426-482.
- Zimmermann, T. M., Clouth, J., Elosge, M., Heurich, M., Schneider, E., Wilhelm, S., e Wolfrath, A. (2013). Patient preferences for outcomes of depression treatment in Germany: A choice-based conjoint analysis study. *Journal of affective disorders*, 148(2), 210-219.

Apêndice I

```
# install.packages("boot")
# install.packages("conjoint")
# install.packages("survival")
# install.packages("moments")

library("boot")
library("conjoint")
library("survival")
library("moments")

#### Banco de dado

# Informações

n = 144; p = 8; r = 3;

# Criando o Design

experiment <- expand.grid(
  Acucar <- c("0% Açúcar", "Com Adoçante"),
  Gordura <- c("0% Gordura", "Baixo Teor"),
  Proteina <- c("Soro do Leite", "Bioativa"))
design <- caFactorialDesign(data = experiment, type = "full")
design

# Replicando o arranjo para "n" indivíduos

design_final <- do.call("rbind", replicate(n, design, simplify = F))
design_final
colnames(design_final) <- c("Acucar", "Gordura", "Proteina")
head(design_final)

# Escolha dos Tratamentos

setwd("c:/dadosR")
escolha <- read.table("escolhas_iogurte.txt", h = T, dec = ",")[,6]

# Impondo o contraste de tratamento

contrasts(design_final$Acucar) <- contr.treatment
contrasts(design_final$Gordura) <- contr.treatment
contrasts(design_final$Proteina) <- contr.treatment

# Banco de dados final

dados <- data.frame(design_final, escolha)
dados

##### Análise Frequentista Tradicional #####

logit = clogit(escolha ~ Acucar + Proteina + Gordura, method = "breslow", data = dados)
summary(logit)
(b <- logit$coef)
(x <- head(model.matrix(logit), p))
(P_y = exp(x%*%b)/sum(exp(x%*%b)))

(ep_logit <- cbind(coef(summary(logit))[, 3]))

#### Intervalos de Confiança baseados na distribuição Z
```

```

IC <- confint(logit)

##### Bootstrap por Pares - Random-X resampling #####

f <- function(formula, data, indices, Xpred)
{
  d <- data[indices,]
  model1 <- clogit(formula, method = "breslow", data = d)
  betas <- model1$coef # Betas
  p <- exp(Xpred%*%betas)/sum(exp(Xpred%*%betas)) # Probabilidades
  re_a <- exp(-betas)[1] # Razão de Escolhas
  re_p <- exp(betas)[2]
  re_g <- exp(-betas)[3]
  return(c(betas, p, re_a, re_g, re_p))
}

result <- boot(data = dados, statistic = f, R = 1999, formula = escolha ~ Acucar + Proteina + Gordura,
Xpred = x)
result

### Médias Bootstrap

(media_boot <- cbind(colMeans(result$t)))

# Intervalos de Confiança Percentil Bootstrap

k = 14
percentil <- matrix(56,14,4)
colnames(percentil) <- c("Q_L", "Q_S", "LI", "LS")

for (i in 1:k)
{
  percentil[i,] <- c(boot.ci(result, conf = 0.95, index = i, typ = "perc")$percent[, -1])
}
percentil

# Gráficos da distribuição empírica dos betas

par(mfrow=c(3,1))
hist(result$t[,1], prob = T, main = "Açúcar", xlab = "", ylab = "Densidade")
abline(v = c(IC[1,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[1,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,2], prob = T, main = "Gordura", xlab = "", ylab = "Densidade")
abline(v = c(IC[2,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[2,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,3], prob = T, main = "Proteína", xlab = "", ylab = "Densidade")
abline(v = c(IC[3,]), col = "blue", lty = 2, lwd = 2)
abline(v = c(percentil[3,c(3,4)]), col = "red", lty = 2, lwd = 2)

# Gráficos da distribuição empírica das Probabilidades de Escolha

par(mfrow=c(4,2))
hist(result$t[,4], prob = T, main = "Prob. Trat. 1", xlab = "", ylab = "Densidade") # p1
abline(v = c(percentil[4,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,5], prob = T, main = "Prob. Trat. 2", xlab = "", ylab = "Densidade") # p2
abline(v = c(percentil[5,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,6], prob = T, main = "Prob. Trat. 3", xlab = "", ylab = "Densidade") # p3
abline(v = c(percentil[6,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,7], prob = T, main = "Prob. Trat. 4", xlab = "", ylab = "Densidade") # p4
abline(v = c(percentil[7,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,8], prob = T, main = "Prob. Trat. 5", xlab = "", ylab = "Densidade") # p5

```

```

abline(v = c(percentil[8,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,9], prob = T, main = "Prob. Trat. 6", xlab = "", ylab = "Densidade") # p6
abline(v = c(percentil[9,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,10], prob = T, main = "Prob. Trat. 7", xlab = "", ylab = "Densidade") # p7
abline(v = c(percentil[10,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,11], prob = T, main = "Prob. Trat. 8", xlab = "", ylab = "Densidade") # p8
abline(v = c(percentil[11,c(3,4)]), col = "red", lty = 2, lwd = 2)

# Gráficos da distribuição empírica da Razão de Escolhas

par(mfrow=c(3,1))
hist(result$t[,12], prob = T, main = "Açúcar", xlab = "", ylab = "Densidade") # re(Acucar)
abline(v = c(percentil[12,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,13], prob = T, main = "Gordura", xlab = "", ylab = "Densidade") # re(Gordura)
abline(v = c(percentil[13,c(3,4)]), col = "red", lty = 2, lwd = 2)
hist(result$t[,14], prob = T, main = "Proteína", xlab = "", ylab = "Densidade") # re(Proteína)
abline(v = c(percentil[14,c(3,4)]), col = "red", lty = 2, lwd = 2)

# Teste de Normalidade Assimetria e Curtose para as Probabilidade

k = 14
teste <- matrix(0,14,8)
colnames(teste) <- c("W", "p", "Ass", "z", "p", "Curt.", "z", "p")
rownames(teste) <- c("b1", "b2", "b3", "p1", "p2", "p3", "p4", "p5", "p6", "p7", "p8", "re1", "re2", "re3")

for (i in 1:k)
{
teste[i,1] <- round(shapiro.test(result$t[,i])$statistic,4)
teste[i,2] <- round(shapiro.test(result$t[,i])$p.value,4)
teste[i,3] <- round(agostino.test(result$t[,i])$statistic[1],4)
teste[i,4] <- round(agostino.test(result$t[,i])$statistic[2],4)
teste[i,5] <- round(agostino.test(result$t[,i])$p.value,4)
teste[i,6] <- round(anscombe.test(result$t[,i])$statistic[1],4)
teste[i,7] <- round(anscombe.test(result$t[,i])$statistic[2],4)
teste[i,8] <- round(anscombe.test(result$t[,i])$p.value,4)
}
teste[4:11,]

# Teste de Normalidade Assimetria e Curtose para a Razão de Escolhas

teste[12:14,]

```

Conclusões Gerais

Com a realização dessa tese pôde-se verificar que o método *Bootstrap* foi uma alternativa metodológica eficiente quando aplicado à *Conjoint Analysis*. A referida técnica produziu estimativas pontuais mais precisas e inferências adicionais que, para serem obtidas sob o ponto de vista clássico ou frequentista, exigiriam do pesquisador conhecimentos específicos de estatística matemática. Logo, a praticidade e o baixo custo computacional do método proposto são suas principais vantagens. Espera-se que o acesso das distribuições empíricas e consequentemente a análise dos intervalos de confiança Percentis *Bootstrap* para as Importâncias Relativas, Probabilidades e Razões de Escolha possam auxiliar gestores e pesquisadores a tomar decisões ou obter conclusões mais confiáveis, baseadas agora em informações probabilísticas.