



UNIVERSIDADE FEDERAL DE PERNAMBUCO
CENTRO DE CIÊNCIAS EXATAS E DA NATUREZA
PROGRAMA DE PÓS-GRADUAÇÃO EM ESTATÍSTICA

JONAS WEVERSON DE ARAÚJO SILVA

INFERÊNCIAS EM MODELOS DE REGRESSÃO SIMPLEX NÃO LINEAR

Recife

2021

JONAS WEVERSON DE ARAÚJO SILVA

INFERÊNCIAS EM MODELOS DE REGRESSÃO SIMPLEX NÃO LINEAR

Tese apresentada ao Programa de Pós-Graduação em Estatística do Centro de Ciências Exatas e da Natureza da Universidade Federal de Pernambuco, como requisito parcial à obtenção do título de doutor em Estatística.

Área de Concentração: Estatística Aplicada

Orientadora: Prof^a. Dr^a. Patrícia Leone Espinheira Ospina

Recife

2021

Catálogo na fonte
Bibliotecário Cristiano Cosme S. dos Anjos, CRB4-2290

S586i Silva, Jonas Weverson de Araújo
 Inferências em modelos de regressão simplex não linear/ Jonas Weverson
 de Araújo silva. – 2021.
 100 f.: il., fig., tab.

 Orientadora: Patrícia Leone Espinheira Ospina.
 Tese (Doutorado) – Universidade Federal de Pernambuco. CCEN,
 Estatística, Recife, 2021.
 Inclui referências e apêndices.

 1. Estatística Aplicada. 2. Bootstrap. 3. Modelo de regressão simplex não
 linear. 4. Pequenas amostras. I. Ospina, Patrícia Leone Espinheira
 (orientadora). II. Título.

 310 CDD (23. ed.) UFPE- CCEN 2021 - 90

JONAS WEVERSON DE ARAÚJO SILVA

INFERÊNCIAS EM MODELOS DE REGRESSÃO SIMPLEX NÃO LINEAR

Tese apresentada ao Programa de Pós-Graduação em Estatística da Universidade Federal de Pernambuco, como requisito parcial para a obtenção do título de Doutor em Estatística.

Aprovada em: 31 de março de 2021.

BANCA EXAMINADORA

Professora Patrícia Leone Espinheira Ospina
UFPE

Professor Francisco Cribari Neto
UFPE

Professor Abraão David Costa do Nascimento
UFPE

Professor Juvêncio Santos Nobre
UFC

Professor Fábio Mariano Bayer
UFSM

Dedico este trabalho a minha mãe, Maria do Livramento, e a meu pai, José Zito.

AGRADECIMENTOS

Primeiramente agradeço a Deus por tudo que Ele tem feito por mim. Por todas as bênçãos, livramentos e por mais esta conquista.

Aos meus pais, Maria do Livramento e José Zito, que sempre me apoiaram e incentivaram a continuar nesta jornada. Por estarem sempre ao meu lado, me dando todo amor, força e orando por mim.

Aos meus irmãos, José Wellington, Joabe Wallison e Wendson que sempre acreditaram em mim.

A Professora Patrícia Espinheira Ospina, minha orientadora, pela paciência, compreensão, pelos enriquecedores conselhos e sugestões, pelas proveitosas discussões e incentivo nas horas de inquietação. Muito obrigado.

Ao meu colega Alisson Silva, por toda ajuda, paciência e por dividir suas ideias e conhecimentos comigo. Sou muito grato a você.

Ao meu colega Fábio Lima, por ajudar no desenvolvimento deste trabalho.

Aos meus amigos de doutorado Pedro Monteiro, Fernanda Clotilde, César Diogo, Adenice Ferreira, Cristina Guedes e Maria Ionerres. Obrigado por estarem sempre dispostos a me ajudar e por terem compartilhado todos os dramas e alegrias dessa trajetória. Em especial, agradeço a Jodavid Ferreira, por sempre me ajudar quando surgia algum problema da área computacional.

A minha amiga Flávia da Silva, por todos os conselhos, encorajamentos e incentivos nos momentos difíceis.

Aos professores da Pós-Graduação, pelos valiosos conhecimentos repassados.

A Valéria e Michelle, pela competência, dedicação e gentileza com os alunos da pós-graduação.

Aos membros da banca examinadora, pelas valiosas sugestões.

A CAPES, pelo apoio financeiro.

Agradeço a todos que contribuíram direta ou indiretamente para a concretização deste sonho.

Muito obrigado.

RESUMO

O objetivo deste trabalho é propor melhoramentos inferenciais baseados em bootstrap (EFRON, 1979) para os parâmetros do modelo de regressão simplex não linear (ESPINHEIRA; SILVA, 2019) em pequenas amostras. Consideramos o estimador de máxima verossimilhança (EMV) tradicional e uma versão do EMV corrigida via bootstrap. Para avaliar a qualidade dos estimadores calculamos o viés relativo e a raiz do erro quadrático médio. Para a estimação intervalar, além do intervalo assintótico baseado na distribuição limite normal do EMV avaliamos também os desempenhos dos intervalos de confiança bootstrap percentil e bootstrap- t . Adicionalmente, investigamos os desempenhos dos testes da razão de verossimilhanças generalizada, escore, Wald, gradiente, a versão corrigida via fator de Bartlett bootstrap e o teste bootstrap duplo rápido com base na estatística da razão de verossimilhanças generalizada. O método bootstrap será uma importante ferramenta para este estudo, visto que através dele será possível contornar diversos problemas de inferências em amostras finitas. O desempenho da inferência por máxima verossimilhança é avaliado numericamente através de simulações de Monte Carlo. Por fim, apresentamos uma aplicação cujo os dados são do departamento de Química da Universidade Nacional da Colômbia (SALAZAR, 2005).

Palavras-chaves: Bootstrap. Modelo de Regressão Simplex não Linear. Pequenas Amostras. Inferência.

ABSTRACT

The objective of this work is to propose inferential improvements based on bootstrap (EFRON, 1979) for the parameters of the nonlinear simplex regression model (ESPINHEIRA; SILVA, 2019) in small samples. We consider the traditional maximum likelihood estimator (MLE) and a version of the MLE corrected via bootstrap. To assess the quality of the estimators, we calculated the relative bias and the root of the mean square error. For interval estimation, in addition to the asymptotic interval based on the normal limit distribution of MLE, we also evaluated the performance of the bootstrap percentile and bootstrap- t confidence intervals. Additionally, investigated the performance of the generalized likelihood ratio, score, Wald, gradient, the version corrected via Bartlett factor bootstrap and the fast double bootstrap test based on generalized likelihood ratio statistics. The bootstrap method will be an important tool for this study, since it will be possible to circumvent several problems of inferences in finite samples. The performance of the inference by maximum likelihood is evaluated numerically through Monte Carlo simulations. Finally, we present an application whose data are from the Chemistry department of the National University of Colombia (SALAZAR, 2005).

Keywords: Bootstrap. Nonlinear Simplex Regression Model. Small Samples. Inference.

LISTA DE FIGURAS

Figura 1 – Curvas de densidade da distribuição simplex variando os valores de μ e σ^2 .	24
Figura 2 – Estimação intervalar para β_2 com $n = 40$, $\mu_t \in (0.02, 0.32)$ e $\lambda \approx 128$.	56
Figura 3 – Estimação intervalar para β_2 com $n = 40$, $\mu_t \in (0.19, 0.86)$ e $\lambda \approx 128$.	56
Figura 4 – Estimação intervalar para β_2 com $n = 40$, $\mu_t \in (0.78, 0.98)$ e $\lambda \approx 128$.	57
Figura 5 – Estimação intervalar para γ_2 com $n = 40$, $\mu_t \in (0.02, 0.32)$ e $\lambda \approx 128$.	57
Figura 6 – Estimação intervalar para γ_2 com $n = 40$, $\mu_t \in (0.19, 0.86)$ e $\lambda \approx 128$.	58
Figura 7 – Estimação intervalar para γ_2 com $n = 40$, $\mu_t \in (0.78, 0.98)$ e $\lambda \approx 128$.	58
Figura 8 – Boxplots dos valores da resposta (a) e das covariáveis vapor d'água (b) e vanádio (c). Dados de craqueamento catalítico fluido (FCC).	60

LISTA DE ALGORITMOS

Algoritmo 1 – Método paramétrico	34
Algoritmo 2 – Intervalo de confiança bootstrap percentil	37
Algoritmo 3 – Intervalo de confiança bootstrap- t	38
Algoritmo 4 – Correção de Bartlett bootstrap	70
Algoritmo 5 – Bootstrap duplo rápido	72

LISTA DE TABELAS

Tabela 1 – Vieses relativos e raízes dos erros quadráticos médios dos EMVs e dos EMVs corrigidos bootstrap (Boot) dos parâmetros do Modelo (3.3) para $\mu_t \approx 0$	42
Tabela 2 – Vieses relativos e raízes dos erros quadráticos médios dos EMVs e dos EMVs corrigidos bootstrap (Boot) dos parâmetros do Modelo (3.3) para $\mu_t \approx 0.5$	43
Tabela 3 – Vieses relativos e raízes dos erros quadráticos médios dos EMVs e dos EMVs corrigidos bootstrap (Boot) dos parâmetros do Modelo (3.3) para $\mu_t \approx 1$	44
Tabela 4 – Taxas de cobertura dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para os parâmetros do modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $1 - \alpha = 0.90$	47
Tabela 5 – Taxas de cobertura dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para os parâmetros do modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $1 - \alpha = 0.95$	48
Tabela 6 – Taxas de cobertura dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para os parâmetros do modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $1 - \alpha = 0.99$	49
Tabela 7 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para β_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $\mu_t \in (0.02, 0.32)$; $\beta_2 = 1.2$; $n = 40$	51

Tabela 8 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para β_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $\mu_t \in (0.02, 0.32)$; $\beta_2 = 1.2$; $n = 80$	52
Tabela 9 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para β_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $\mu_t \in (0.02, 0.32)$; $\beta_2 = 1.2$; $n = 120$	53
Tabela 10 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para β_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $n = 40$; $\lambda \approx 128$	54
Tabela 11 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para γ_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $n = 40$; $\lambda \approx 128$	55
Tabela 12 – Estimativas de máxima verossimilhança $\hat{\theta}$, estimativas com viés corrigido via bootstrap ($\bar{\theta}$), erros-padrão (EP) e p -valores associados aos testes Z para os parâmetros do Modelo 3.4. Dados de craqueamento catalítico fluido (FCC).	61
Tabela 13 – Estimativas intervalares assintótica (ICA), bootstrap-t (ICBT) e percentil (ICPE) para os parâmetros do Modelo 3.4. Dados de craqueamento catalítico fluido (FCC)	63

Tabela 14 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_3 = \beta_4 = 0$ ($r = 2$). $\lambda \approx 9$	75
Tabela 15 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_3 = \beta_4 = 0$ ($r = 2$). $\lambda \approx 53$	76
Tabela 16 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_3 = \beta_4 = 0$ ($r = 2$). $\lambda \approx 101$	77
Tabela 17 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_2 = 1$ ($r = 1$). $\sigma^2 = 0.5$	79
Tabela 18 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_2 = 1$ ($r = 1$). $\sigma^2 = 1.5$	80
Tabela 19 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_2 = 1$ ($r = 1$). $\sigma^2 = 3.0$	81

Tabela 20 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_2 = 1$ ($r = 1$). $\sigma^2 = 6.0$	82
Tabela 21 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \gamma_2 = \gamma_3 = 0$ ($r = 2$). $\sigma^2 = 0.5$	83
Tabela 22 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \gamma_2 = \gamma_3 = 0$ ($r = 2$). $\sigma^2 = 1.5$	84
Tabela 23 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \gamma_2 = \gamma_3 = 0$ ($r = 2$). $\sigma^2 = 3.0$	85
Tabela 24 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \gamma_2 = \gamma_3 = 0$ ($r = 2$). $\sigma^2 = 6.0$	86
Tabela 25 – Estimativas de máxima verossimilhança $\hat{\theta}$ e erros-padrão (EP) para os parâmetros do modelo $M1$. Dados de craqueamento catalítico fluido (FCC). .	89
Tabela 26 – p -valores dos testes de $H_0 : \beta_2 = 0$	90
Tabela 27 – p -valores dos testes de $H_0 : \gamma_1 = \gamma_2 = 0$	90
Tabela 28 – p -valores dos testes individuais sobre os parâmetros do modelo $M2$	90
Tabela 29 – Critérios de seleção para os modelos $M2$ e $M3$	91
Tabela 30 – Estimativas de máxima verossimilhança $\hat{\theta}$ e erros-padrão (EP) para os parâmetros do modelo $M3$. Dados de craqueamento catalítico fluido (FCC). .	91

SUMÁRIO

1	INTRODUÇÃO	16
1.1	ORGANIZAÇÃO DA TESE	20
1.2	SUPORTE COMPUTACIONAL	20
2	DISTRIBUIÇÃO SIMPLEX	22
2.1	PROPRIEDADES DA DISTRIBUIÇÃO SIMPLEX	23
2.2	MODELO DE REGRESSÃO SIMPLEX NÃO LINEAR	24
3	ESTIMAÇÃO PONTUAL E INTERVALAR	30
3.1	INTRODUÇÃO	30
3.2	ESTIMAÇÃO PONTUAL DOS PARÂMETROS DO MODELO	30
3.3	MÉTODO BOOTSTRAP	32
3.4	CORREÇÃO DE VIÉS DOS ESTIMADORES PONTUAIS	34
3.5	ESTIMAÇÃO INTERVALAR DOS PARÂMETROS DO MODELO	34
3.5.1	Intervalo de Confiança Bootstrap Percentil	36
3.5.2	Intervalo de Confiança Bootstrap-t	37
3.6	RESULTADOS NUMÉRICOS	39
3.6.1	Resultados da Avaliação Numérica da Estimação Pontual	40
3.6.2	Resultados Numéricos sobre Intervalos de Confiança	45
3.7	APLICAÇÃO: DADOS DE CRAQUEAMENTO CATALÍTICO FLUIDO (FCC)	59
4	TESTES DE HIPÓTESES	64
4.1	INTRODUÇÃO	64
4.2	TESTES ASSINTÓTICOS	64
4.2.1	Teste da Razão de Verossimilhanças Generalizada	65
4.2.2	Teste Escore de Rao	66
4.2.3	Teste de Wald	66
4.2.4	Teste Gradiente	66
4.3	TESTES DE HIPÓTESES EM REGRESSÃO SIMPLEX NÃO LINEAR SOBRE O VETOR DE PARÂMETROS β	67
4.4	TESTES DE HIPÓTESES EM REGRESSÃO SIMPLEX NÃO LINEAR SOBRE O VETOR DE PARÂMETROS γ	68
4.5	TESTES DE HIPÓTESES BASEADOS EM BOOTSTRAP	69

4.5.1	Correção de Bartlett Bootstrap	69
4.5.2	Teste Bootstrap Duplo Rápido	70
4.6	RESULTADOS NUMÉRICOS	72
4.6.1	Cenário 1	73
4.6.2	Cenário 2	76
4.6.3	Cenário 3	79
4.7	CRITÉRIOS PARA A SELEÇÃO DE MODELOS	87
4.8	APLICAÇÃO: DADOS DE CRAQUEAMENTO CATALÍTICO FLUIDO (FCC)	88
5	CONSIDERAÇÕES FINAIS	92
5.1	TRABALHOS FUTUROS	93
	REFERÊNCIAS	94
	APÊNDICE A – FUNÇÃO ESCORE PARA O PARÂMETRO μ	98
	APÊNDICE B – CONDIÇÕES DE REGULARIDADE (C.R.)	99

1 INTRODUÇÃO

Os modelos de regressão lineares são amplamente utilizados nas mais diversas áreas de conhecimento. Atualmente vêm surgindo várias propostas de modelos de regressão para variáveis resposta duplamente limitadas, que assumem valores contínuos em (a, b) , em que a e b são conhecidos e $-\infty < a < b < \infty$, assim, tal suporte pode ser facilmente transformado para o intervalo unitário.

Neste contexto, em que $y \in (0, 1)$ ($y \in (a, b)$), o modelo linear normal é inadequado, pois além da possibilidade de ocorrência de valores ajustados menores que $0(a)$ ou maiores que $1(b)$, em geral, os dados apresentam assimetria e heteroscedasticidade, violando as suposições usuais de tal modelo. Desta forma, parece mais apropriado considerar modelos baseados em distribuições com suporte naturalmente no $(0, 1)$ como é o caso do modelo de regressão simplex proposto por Barndorff-Nielsen e Jørgensen (1991), por exemplo.

A distribuição simplex foi desenvolvida a partir da distribuição gaussiana inversa generalizada e faz parte da classe dos modelos de dispersão definidos por Jørgensen (1997), que estendem os modelos lineares generalizados (MCCULLAGH; NELDER, 1989). Vários trabalhos foram realizados utilizando esta distribuição. Por exemplo, Song e Tan (2000) utilizaram essa distribuição para avaliar dados longitudinais considerando o parâmetro de dispersão constante, utilizando equações de estimação generalizada. Song, Qiu e Tan (2004) modificaram essa abordagem com a suposição de que o parâmetro de dispersão varia ao longo das observações. Venezuela (2008) desenvolveu equações de estimação para o modelo de regressão simplex com medidas repetidas e apresentou algumas medidas de diagnóstico. Miyashiro (2008) propôs o modelo de regressão simplex com dispersão constante, apresentando várias propriedades e medidas de diagnóstico. Com base nos modelos de dispersão, Song (2009) apresenta uma revisão teórica de análise de regressão e, através do método de máxima verossimilhança, estima os parâmetros do modelo. Santos (2011) apresenta algumas correções para o resíduo de Pearson do modelo simplex com dispersão constante. Utilizando o enfoque bayesiano e simulações de Monte Carlo, López (2013) avalia os estimadores dos parâmetros do modelo simplex com dispersão variável. Silva (2015) apresenta duas novas medidas de diagnóstico para o modelo de regressão simplex, considerando os modelos com dispersão constante e dispersão variável. Espinheira e Silva (2019) propuseram a classe de modelos de regressão simplex não linear, em que estimam os parâmetros do modelo através do método de máxima verossimilhança

e derivaram as quantidades de influência local. Silva (2019) propôs um novo resíduo para a classe de modelos de regressão simplex, e versões das estatísticas PRESS e dos coeficientes de predição P^2 para o modelo de regressão simplex linear e não linear. Recentemente, Liu et al. (2020) apresentaram a distribuição simplex inflacionada em zero e um para modelar dados de proporção. Os autores introduziram um novo algoritmo para calcular as estimativas de máxima verossimilhança dos parâmetros da distribuição simplex sem covariáveis, além disso, desenvolveram métodos de inferência baseados em verossimilhança para o modelo de regressão usando esta nova distribuição.

Outras abordagens para a modelagem de dados limitados são os modelos de regressão beta (FERRARI; CRIBARI-NETO, 2004), Kumaraswamy (MITNIK; BAEK, 2013), Johnson S_B (LEMONTE; BAZAN, 2016), gama unitária (MOUSA; EL-SHEIKH; ABDEL-FATTAH, 2016). Artigos recentemente publicados mostram eventuais vantagens do uso desta última distribuição em relação à distribuição beta (GUEDES; CRIBARI-NETO; ESPINHEIRA, 2020; ROCHA; ESPINHEIRA; CRIBARI-NETO, 2021). Além desses modelos, temos o modelo de regressão CDF-quantile (SMITHSON; SHOU, 2017), no qual é baseado em uma família de distribuições de dois parâmetros obtida pela aplicação da função de distribuição cumulativa à função quantil de outra distribuição. Recentemente, (SILVA, 2019) mostrou que quando os dados estão concentrados nos extremos do intervalo unitário padrão, o processo de estimação por máxima verossimilhança do modelo simplex é mais robusto que o do modelo de regressão beta.

O estudo do comportamento de estimadores de máxima verossimilhança (EMVs) em pequenas amostras é uma importante área de pesquisa. Estes estimadores podem ser viesados quando o tamanho da amostra é pequeno ou até mesmo moderado. O viés é de fato uma medida de risco médio. A média do risco em substituir o verdadeiro valor do parâmetro por um valor estimado plausível. O viés também pode ser visto o quão distante em média está um estimador do verdadeiro valor do parâmetro. Assim, é desejável a obtenção de estimadores com viés reduzido em amostras finitas. Quando o tamanho de amostra é grande o viés tende a zero. Na literatura, existem diversas maneiras de obter estimadores menos viesados em pequenas amostras. Aqui, faremos uso de uma correção de viés obtida a partir do método bootstrap (EFRON, 1979).

Na inferência estatística é de fundamental importância associar confiabilidade às estimativas pontuais do modelo, e uma das formas de fazer isto é através da construção das estimativas intervalares dos parâmetros. Os intervalos de confiança podem ser obtidos através da distribuição assintótica dos estimadores de máxima verossimilhança, que podem requerer grandes

amostras para garantir a validade das aproximações para as distribuições. Em pequenas amostras, uma alternativa para a construção de intervalos de confiança com bom desempenho tanto em relação à taxa de cobertura do verdadeiro valor do parâmetro quanto ao comprimento do intervalo, é o método bootstrap. Aqui, vamos considerar os intervalos de confiança bootstrap percentil e bootstrap- t (EFRON; TIBSHIRANI, 1994), nos quais apresentam níveis de cobertura empíricos mais próximos da verdadeira probabilidade de cobertura.

Outro aspecto importante na inferência estatística são os testes de hipóteses, que têm por objetivo tomar decisões sobre os parâmetros do modelo com uma probabilidade de erro pré-estabelecida. Inferências em grandes amostras nos modelos de regressão podem ser realizadas através dos testes da razão de verossimilhanças generalizada (WILKS, 1938), escore (RAO, 1948), Wald (WALD, 1943) e gradiente (TERRELL, 2002). As estatísticas destes testes baseiam-se na função de verossimilhança, que apresenta várias propriedades de otimalidade. Sob condições usuais de regularidade (ver Apêndice B) e sob a hipótese nula, as distribuições limite destas estatísticas são qui-quadrado, ou seja, as estatísticas são assintoticamente equivalentes. Estes testes são realizados a partir de valores críticos válidos para grandes amostras, o que pode causar considerável distorção no tamanho do teste em amostras pequenas ou moderadas. Esta distorção pode ser contornada através de correções entre as quais baseadas no método bootstrap. Na literatura, existem vários testes utilizando este método. Aqui, faremos uso do teste da razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (ROCKE, 1989) e da razão de verossimilhanças generalizada bootstrap duplo rápido (DAVIDSON; MACKINNON, 2000). Utilizando estes testes é possível obter uma menor distorção de tamanho comparativamente aos testes assintóticos usuais.

No contexto de modelagem de dados contínuos limitados, vários autores já realizaram estudos sobre inferência baseada no método de estimação por máxima verossimilhança. Simas, Barreto-Souza e Rocha (2010) propõem o modelo de regressão beta não linear e apresentam estudos e propostas para os estimadores de máxima verossimilhança do modelo. Ferrari e Pinheiro (2011) apresentam correções da estatística da razão de verossimilhanças generalizada baseadas em Skovgaard (2001) para a classe de modelos de regressão beta. Bayer e Cribari-Neto (2013) desenvolvem testes da razão de verossimilhanças generalizada (RV) corrigidos por Bartlett para essa mesma classe de modelos. Cribari-Neto e Lima (2019) propõem um novo intervalo de predição baseado em bootstrap para respostas ausentes em regressão beta. Os autores também avaliam o impacto da especificação incorreta do modelo nas coberturas empíricas de diferentes intervalos de predição, e investigam o impacto da especificação incorreta

do modelo em três intervalos de predição bootstrap. Lima e Cribari-Neto (2020) abordam a inferência de teste em pequenas amostras na classe de modelos de regressão beta. Os autores consideram o teste RV e suas versões bootstrap, mostram que o teste RV padrão tende a ser bastante liberal em pequenas amostras e que os testes baseados em bootstrap fornecem inferências mais confiáveis, mesmo quando o tamanho da amostra é muito pequeno. Guedes, Cribari-Neto e Espinheira (2020) utilizaram a mesma estratégia apresentada por Ferrari e Pinheiro (2011) para a correção do teste RV , porém considerando à distribuição gama unitária.

Neste trabalho, abordaremos a classe de modelos de regressão simplex proposta por Espinheira e Silva (2019) em que a média da variável resposta e o parâmetro de dispersão estão relacionados a covariáveis por meio de preditores não lineares. Temos por objetivo propor melhoramentos inferenciais baseados em bootstrap para os parâmetros que indexam o modelo em pequenas amostras. Consideramos o EMV tradicional e uma versão do EMV corrigida via bootstrap. Para avaliar a qualidade dos estimadores consideramos o viés relativo e a raiz do erro quadrático médio. Será considerada a estimação intervalar dos parâmetros através da construção de intervalos de confiança do tipo assintótico, bootstrap percentil e bootstrap- t . Avaliamos diversos aspectos da estimação intervalar, a saber: a cobertura empírica, o limite inferior (superior) médio, o comprimento médio e o balanceamento. Avaliamos os desempenhos dos tamanhos dos testes da razão de verossimilhanças generalizada, score, Wald e gradiente, a versão corrigida via fator de Bartlett bootstrap e o teste bootstrap duplo rápido com base na estatística da razão de verossimilhanças generalizada. Vale destacar que o método bootstrap será uma importante ferramenta para este estudo, visto que através dele podemos contornar diversos problemas de inferências em amostras finitas. O desempenho da inferência por MV é avaliado numericamente através de simulações de Monte Carlo. Por fim, apresentamos uma aplicação cujo os dados são do departamento de Química da Universidade Nacional da Colômbia (SALAZAR, 2005) e composto por 28 observações. A aplicação consiste no processo de craqueamento catalítico fluido que é usado para converter hidrocarbonetos de alto peso molecular em pequenas moléculas de maior valor comercial, através do contato destes com um catalisador. Segundo Salazar (2005), o principal catalizador do processo é o zeólito Y . Outra substância importante que participa do processo de catalização é o Vanádio. O processo também depende da temperatura. O interesse é modelar a porcentagem de cristalinidade do zeólito Y com base em diferentes concentrações de vanádio e vapor, considerando dois valores para a temperatura do processo (SALAZAR, 2005).

1.1 ORGANIZAÇÃO DA TESE

O presente trabalho se encontra dividido em cinco capítulos. Neste capítulo, abordamos de forma geral alguns temas que serão tratados neste trabalho, descrevemos os objetivos e apresentamos informações sobre o suporte computacional utilizado.

No segundo capítulo apresentamos a classe de modelos de dispersão com as suas respectivas propriedades, o modelo de regressão simplex não linear e as expressões analíticas para o vetor escore, matriz de informação de Fisher e sua inversa.

No terceiro capítulo discutimos a estimação pontual e intervalar dos parâmetros do modelo por máxima verossimilhança. Apresentamos o método bootstrap, o EMV corrigido via bootstrap e os intervalos de confiança bootstrap que utilizamos, a saber: os intervalos de confiança bootstrap percentil e bootstrap-t. Além disso, com base em simulações de Monte Carlo para o modelo de regressão simplex não linear investigamos o comportamento dos estimadores de máxima verossimilhança e sua versão bootstrap. Adicionalmente, analisamos as diferentes estimativas intervalares dos parâmetros deste modelo. Por fim, aplicamos o modelo de regressão simplex não linear a um conjunto de dados reais.

No quarto capítulo abordamos os testes da razão de verossimilhanças generalizada, escore, Wald, gradiente e os testes baseados em bootstrap como o da razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap e da razão de verossimilhanças generalizada bootstrap duplo rápido. Além disso, avaliamos através de simulações de Monte Carlo o comportamento dos testes abordados e aplicamos os testes a um conjunto de dados reais. Finalmente, o quinto capítulo contém as principais conclusões deste trabalho.

1.2 SUPORTE COMPUTACIONAL

As avaliações numéricas deste trabalho foram realizadas utilizando a linguagem de programação matricial Ox para os sistemas operacionais Linux e Windows. Esta linguagem foi criada e desenvolvida por Jurgen Doornik em 1994 na Universidade de Oxford (Inglaterra), é distribuída gratuitamente para uso acadêmico e se encontra disponível em <http://www.doornik.com>. Para mais informações sobre o Ox, ver Doornik (2001), Doornik (2006) e Doornik (2009).

Os resultados gráficos apresentados neste trabalho foram obtidos utilizando o ambiente de programação R para o sistema operacional Windows. O R foi criado por Ross Ihaka e

Robert Gentleman na Universidade de Auckland e é um sistema integrado para computação estatística e geração de gráficos que se encontra disponível gratuitamente através do site <http://www.R-project.org>. Para mais detalhes, ver Cribari-Neto e Zarkos (1999), Dalgaard (2002) e Venables e Ripley (2002).

Esta tese foi digitada utilizando o sistema de tipografia $\text{\LaTeX}$ desenvolvido por Lamport em 1985, que consiste em uma série de macros ou rotinas do sistema $\text{\TeX}$. Criado por Donald Knuth na Universidade de Stanford, tem por objetivo facilitar o desenvolvimento da edição de textos (KNUTH, 1986). Detalhes sobre o sistema de tipografia $\text{\LaTeX}$ podem ser encontrados em Lamport (1994) ou através do site <http://www.tex.ac.uk/CTAN/latex>.

2 DISTRIBUIÇÃO SIMPLEX

A distribuição simplex é utilizada para modelar dados restritos ao intervalo $(0,1)$. Esta distribuição proposta por Barndorff-Nielsen e Jørgensen (1991) pertence à classe de modelos de dispersão introduzidos por Jørgensen (1997), que estendem os modelos lineares generalizados (NELDER; WEDDERBURN, 1972). Dizemos que uma variável aleatória y possui distribuição pertencente à classe de modelos de dispersão, denotada por $y \sim DM(\mu, \sigma^2)$, se sua função densidade de probabilidade é da forma

$$p(y; \mu, \sigma^2) = a(y; \sigma^2) \exp \left\{ -\frac{1}{2\sigma^2} d(y; \mu) \right\}, \quad y \in C, \quad (2.1)$$

em que C é o suporte da distribuição, $\mu \in \Omega$ é o parâmetro de posição/localização, $\sigma^2 > 0$ é o parâmetro de dispersão e $a(\cdot) \geq 0$ é uma função regularmente adequada em termos de ser funcionalmente independente de μ , Ω sendo o espaço paramétrico de μ . A função $d(y; \mu)$ é chamada componente do desvio definido em $(y, \mu) \in (C, \Omega)$ se satisfaz

$$d(y; \mu) = 0, \forall y = \mu \in \Omega \quad \text{e} \quad d(y; \mu) > 0, \forall y \neq \mu.$$

A função de variância para os modelos de dispersão é uma função $V : \Omega \rightarrow (0, \infty)$ definida por

$$V(\mu) = \frac{2}{\frac{\partial^2 d(y; \mu)}{\partial \mu^2} \Big|_{y=\mu}},$$

em que a função $d(y; \mu)$ é contínua e duas vezes diferenciável com respeito a μ , com

$$\frac{\partial^2 d(y; \mu)}{\partial \mu^2} \Big|_{y=\mu} > 0, \quad \forall \mu \in \Omega.$$

Na literatura existem várias distribuições discretas e contínuas que pertencem a classe de modelos de dispersão, entre as quais podemos citar as distribuições: normal, normal inversa, gama, von Mises, Poisson, Binomial, Binomial negativa, entre outras. Em particular, se uma variável aleatória y segue a distribuição simplex denotada por $S^-(\mu, \sigma^2)$ com parâmetros $0 < \mu < 1$ e $\sigma^2 > 0$, a expressão (2.1) toma a forma

$$p(y; \mu, \lambda) = [2\pi\sigma^2 \{y(1-y)\}^3]^{-1/2} \exp \left\{ -\frac{1}{2\sigma^2} d(y; \mu) \right\}, \quad y \in (0, 1), \quad (2.2)$$

em que o componente do desvio $d(y; \mu)$ é dado por

$$d(y; \mu) = \frac{(y - \mu)^2}{y(1-y)\mu^2(1-\mu)^2}.$$

A função de variância para a distribuição simplex é expressa como $V(\mu) = \mu^3(1 - \mu)^3$. A esperança e variância dessa distribuição são dadas, respectivamente, por

$$\mathbb{E}(y) = \mu$$

e

$$\text{Var}(y) = \mu(1 - \mu) - \sqrt{\frac{1}{2\sigma^2}} \exp \left\{ \frac{1}{\sigma^2 \mu^2 (1 - \mu)^2} \right\} \Gamma \left(\frac{1}{2}, \frac{1}{2\sigma^2 \mu^2 (1 - \mu)^2} \right),$$

em que $\Gamma(a, b)$ é a função gama incompleta, definida por $\Gamma(a, b) = \int_b^\infty x^{a-1} e^{-x} dx$. Para mais detalhes sobre essas propriedades, ver Jørgensen (1997). A distribuição simplex é bastante flexível para modelar dados no intervalo contínuo $(0, 1)$, apresentando diferentes formas de acordo com os valores dos parâmetros que indexam a distribuição. Na Figura 1 apresentamos alguns exemplos, como a forma de J para $S^-(0.9, 36)$, a forma de U para $S^-(0.5, 121)$ e a forma de J invertido para $S^-(0.1, 36)$, além das formas comuns, a saber: assimétrica à esquerda, assimétrica à direita e simétrica. Além disso, ao contrário da distribuição beta, o modelo simplex é muito útil para acomodar dados com distribuições bimodais (Figura 1, $S^-(0.5, 20)$).

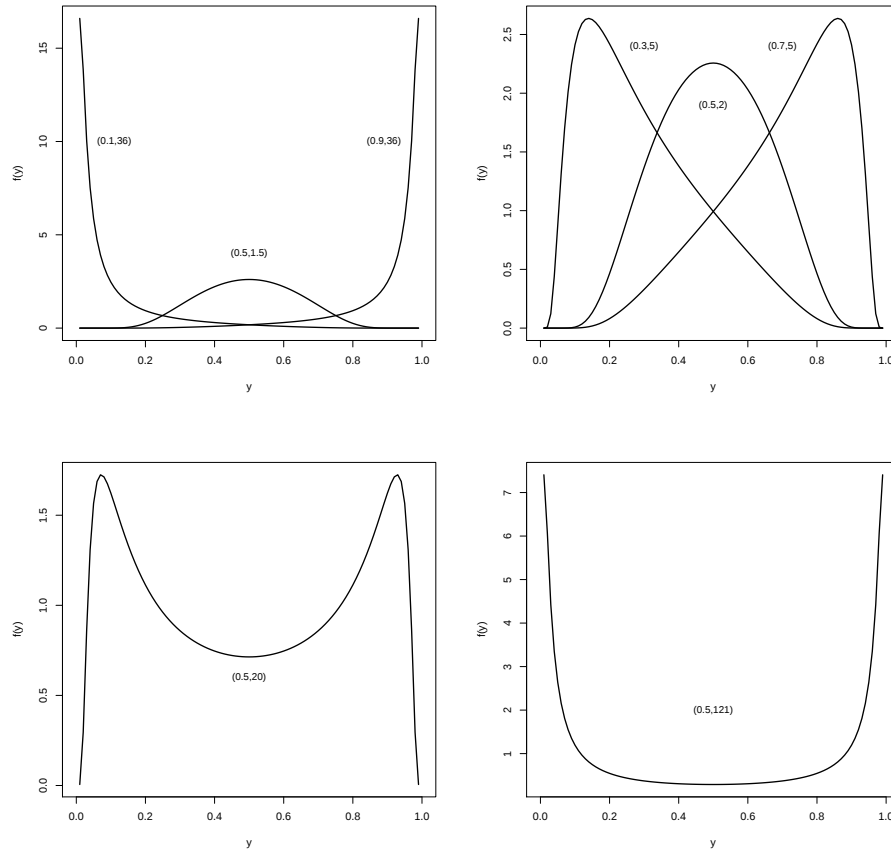
2.1 PROPRIEDADES DA DISTRIBUIÇÃO SIMPLEX

A seguir, apresentamos algumas propriedades da distribuição simplex que serão utilizadas no decorrer deste trabalho.

Proposição 2.1: Seja y uma variável aleatória tal que $y \sim S^-(\mu, \sigma^2)$. Então, valem as seguintes propriedades:

- (i) $\mathbb{E}[d(y; \mu)] = \sigma^2$;
- (ii) $\mathbb{E}[(y - \mu)d'(y; \mu)] = -2\sigma^2$;
- (iii) $\mathbb{E}[(y - \mu)d(y; \mu)] = 0$;
- (iv) $\mathbb{E}[(y - \mu)d''(y; \mu)] = 0$;
- (v) $\frac{1}{2}\mathbb{E}[d''(y; \mu)] = \frac{3\sigma^2}{\mu(1-\mu)} + \frac{1}{\mu^3(1-\mu)^3}$;
- (vi) $\text{Var}[d(y; \mu)] = 2(\sigma^2)^2$;
- (vii) $\mathbb{E}[d'(y; \mu)] = 0$,

Figura 1 – Curvas de densidade da distribuição simplex variando os valores de μ e σ^2 .



Fonte: próprio autor.

em que $d'(y; \mu) = \partial d(y; \mu) / \partial \mu$ e $d''(y; \mu) = \partial^2 d(y; \mu) / \partial \mu^2$. As provas dos itens (i), (ii) e (iii) podem ser encontradas em Song e Tan (2000). Já as provas dos itens (iv), (v) e (vi) se encontram em Silva (2016). Por fim, a prova do item (vii) está em Song, Qiu e Tan (2004).

2.2 MODELO DE REGRESSÃO SIMPLEX NÃO LINEAR

Sejam $y_1, \dots, y_n$ variáveis aleatórias independentes, em que cada y_t , $t = 1, \dots, n$, segue uma distribuição simplex cuja função densidade de probabilidade é dada por (2.2) com média μ_t e parâmetro de dispersão σ_t^2 .

O modelo de regressão simplex não linear proposto por Espinheira e Silva (2019) é definido por (2.2) e pelas componentes sistemáticas

$$g(\mu_t) = f_1(x_t^\top; \beta) = \eta_t \quad \text{e} \quad h(\sigma_t^2) = f_2(z_t^\top; \gamma) = \zeta_t, \quad t = 1, \dots, n, \quad (2.3)$$

em que $\beta = (\beta_1, \dots, \beta_k)^\top$ e $\gamma = (\gamma_1, \dots, \gamma_q)^\top$ são vetores de parâmetros de regressão desconhecidos tais que $\beta \in \mathbb{R}^k$ e $\gamma \in \mathbb{R}^q$, $k+q < n$, $\eta_t = (\eta_1, \dots, \eta_n)^\top$ e $\zeta_t = (\zeta_1, \dots, \zeta_n)^\top$ são pre-

ditores não lineares e $x_t^\top = (x_{t1}, \dots, x_{tk_1})$ e $z_t^\top = (z_{t1}, \dots, z_{tq_1})$ são, respectivamente, k_1 e q_1 observações de covariáveis conhecidas, que podem coincidir total ou parcialmente de tal modo que $k_1 \leq k$ e $q_1 \leq q$. Nos modelos lineares $k_1 = k$, $q_1 = q$, e portanto, $x_t^\top = (x_{t1}, \dots, x_{tk_1})$ e $z_t^\top = (z_{t1}, \dots, z_{tq_1})$ são, respectivamente, as t -ésimas linhas das matrizes X e Z , para $t = 1, \dots, n$. Os modelos lineares são um caso particular dos modelos não lineares. Quando temos a não linearidade nos parâmetros, pelo menos um dos $\partial f_1(\cdot; \beta)/\partial \beta_j, j = 1, \dots, k$ depende de $\beta^\top = (\beta_1, \dots, \beta_k)$ e, pelo menos um dos $\partial f_2(\cdot; \gamma)/\partial \gamma_l, l = 1, \dots, q$ depende de $\gamma^\top = (\gamma_1, \dots, \gamma_q)$. Para modelos simplex lineares, essas derivadas dependem apenas das covariáveis $x_1, \dots, x_k$ e $z_1, \dots, z_q$, respectivamente, e assim $\tilde{\mathcal{X}} = \partial \eta / \partial \beta = X$ e $\tilde{\mathcal{Z}} = \partial \zeta / \partial \gamma = Z$. Neste trabalho, vamos considerar o caso de não linearidade.

As funções de ligação $g : (0, 1) \rightarrow \mathbb{R}$ e $h : (0, \infty) \rightarrow \mathbb{R}$ são estritamente monótonas e ao menos duas vezes diferenciáveis. Diferentes funções de ligação podem ser escolhidas para g e h . Por exemplo, para μ podemos usar a função logito $g(\mu) = \log\{\mu/(1 - \mu)\}$, a função probito $g(\mu) = \Phi^{-1}(\mu)$, em que $\Phi(\cdot)$ denota a função de distribuição normal padrão, a função log-log $g(\mu) = \log\{-\log(\mu)\}$ e a função complemento log-log $g(\mu) = \log\{-\log(1 - \mu)\}$, entre outras. Como $\sigma^2 > 0$, podemos usar a função logarítmica $h(\sigma^2) = \log(\sigma^2)$ e a função identidade $h(\sigma^2) = \sigma^2$, com atenção especial para a positividade das estimativas. Para mais detalhes, ver Atkinson (1985) e McCullagh e Nelder (1989). Em (2.3) temos que $g(\mu_t)$ e $h(\sigma_t^2)$, $t = 1, \dots, n$, representam os submodelos da média e da dispersão, respectivamente.

A seguir, apresentaremos a função escore, a matriz de informação de Fisher e sua inversa para o vetor de parâmetros $\theta = (\beta^\top, \gamma^\top)^\top$. Esses resultados foram obtidos por Espinheira e Silva (2019). Eles serão importantes para a obtenção dos estimadores de máxima verossimilhança, construção de intervalos de confiança assintóticos e para a realização de testes de hipóteses sobre os parâmetros do modelo.

Baseado em (2.2) e uma amostra de n observações independentes, temos que o logaritmo da função de verossimilhança é dado por

$$\ell(\beta, \gamma) = \sum_{t=1}^n \ell_t(\mu_t, \sigma_t^2), \quad (2.4)$$

em que

$$\ell_t(\mu_t, \sigma_t^2) = -\frac{1}{2} \log 2\pi - \frac{1}{2} \log \sigma_t^2 - \frac{3}{2} \log\{y_t(1 - y_t)\} - \frac{1}{2\sigma_t^2} d(y_t; \mu_t).$$

Os componentes do vetor escore $U_\beta(\beta, \gamma)$ são obtidos pela diferenciação do logaritmo da

função de verossimilhança (2.4) com respeito a β_i , $i = 1, \dots, k$. Assim, temos que

$$\begin{aligned} \frac{\partial \ell(\beta, \gamma)}{\partial \beta_i} &= \sum_{t=1}^n \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \beta_i} \\ &= \sum_{t=1}^n \left\{ \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t} \frac{d\mu_t}{d\eta_t} \frac{\partial \eta_t}{\partial \beta_i} \right\} \\ &= \sum_{t=1}^n \left\{ \frac{1}{\sigma_t^2} \frac{(y_t - \mu_t)}{\mu_t(1 - \mu_t)} \left[d(y_t; \mu_t) + \frac{1}{\mu_t^2(1 - \mu_t)^2} \right] \frac{1}{g'(\mu_t)} \frac{\partial \eta_t}{\partial \beta_i} \right\}, \end{aligned}$$

em que $d\mu_t/d\eta_t = 1/g'(\mu_t)$ e

$$\begin{aligned} \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t} &= -\frac{1}{2\sigma_t^2} d'(y_t; \mu_t) \\ &= \frac{1}{\sigma_t^2} \frac{(y_t - \mu_t)}{\mu_t(1 - \mu_t)} \left[d(y_t; \mu_t) + \frac{1}{\mu_t^2(1 - \mu_t)^2} \right]. \end{aligned} \quad (2.5)$$

(Para mais detalhes, ver Apêndice A). Matricialmente, a função escore para o vetor $\beta = (\beta_1, \dots, \beta_k)^\top$ pode ser expressa da seguinte forma:

$$U_\beta(\beta, \gamma) = \widetilde{\mathcal{X}}^\top SUT(y - \mu),$$

em que $\widetilde{\mathcal{X}} = \frac{\partial \eta}{\partial \beta}$ é uma matriz de derivadas de dimensão $n \times k$, $y = (y_1, \dots, y_n)^\top$ e $\mu = (\mu_1, \dots, \mu_n)^\top$ são vetores e $U = \text{diag}\{u_1, \dots, u_n\}$ é a matriz diagonal, em que

$$u_t = \frac{1}{\mu_t(1 - \mu_t)} \left\{ d(y_t; \mu_t) + \frac{1}{\mu_t^2(1 - \mu_t)^2} \right\}, \quad t = 1, \dots, n. \quad (2.6)$$

Além disso,

$$S = \text{diag} \left\{ \frac{1}{\sigma_1^2}, \dots, \frac{1}{\sigma_n^2} \right\} \quad \text{e} \quad T = \text{diag} \left\{ \frac{1}{g'(\mu_1)}, \dots, \frac{1}{g'(\mu_n)} \right\}. \quad (2.7)$$

Agora vamos obter os componentes do vetor escore $U_\gamma(\beta, \gamma)$ diferenciando o logaritmo da função de verossimilhança (2.4) com respeito a γ_j , $j = 1, \dots, q$, que são dados por

$$\begin{aligned} \frac{\partial \ell(\beta, \gamma)}{\partial \gamma_j} &= \sum_{t=1}^n \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \gamma_j} \\ &= \sum_{t=1}^n \left\{ \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \sigma_t^2} \frac{d\sigma_t^2}{d\zeta_t} \frac{\partial \zeta_t}{\partial \gamma_j} \right\} \\ &= \sum_{t=1}^n \left\{ \left[-\frac{1}{2\sigma_t^2} + \frac{d(y_t; \mu_t)}{2(\sigma_t^2)^2} \right] \frac{1}{h'(\sigma_t^2)} \frac{\partial \zeta_t}{\partial \gamma_j} \right\}, \end{aligned}$$

em que $d\sigma_t^2/d\zeta_t = 1/h'(\sigma_t^2)$ e

$$\frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \sigma_t^2} = -\frac{1}{2\sigma_t^2} + \frac{d(y_t; \mu_t)}{2(\sigma_t^2)^2}. \quad (2.8)$$

Matricialmente, a função escore para o vetor $\gamma = (\gamma_1, \dots, \gamma_q)^\top$ pode ser escrita como

$$U_\gamma(\beta, \gamma) = \tilde{Z}^\top H a,$$

em que $\tilde{Z} = \frac{\partial \zeta}{\partial \gamma}$ é uma matriz de derivadas de dimensão $n \times q$, $a = (a_1, \dots, a_n)^\top$ com

$$a_t = -\frac{1}{2\sigma_t^2} + \frac{d(y_t; \mu_t)}{2(\sigma_t^2)^2} \quad \text{e} \quad H = \text{diag} \left\{ \frac{1}{h'(\sigma_1^2)}, \dots, \frac{1}{h'(\sigma_n^2)} \right\}. \quad (2.9)$$

Portanto, os componentes do vetor escore $(U_\beta(\beta, \gamma)^\top, U_\gamma(\beta, \gamma)^\top)^\top$ são dados por

$$U_\beta(\beta, \gamma) = \tilde{\mathcal{X}}^\top SUT(y - \mu) \quad \text{e} \quad U_\gamma(\beta, \gamma) = \tilde{Z}^\top H a.$$

Para obtermos a matriz de informação de Fisher para os vetores de parâmetros β e γ , devemos primeiro calcular as derivadas de segunda ordem do logaritmo da função de verossimilhança (2.4). Assim, para $i, p = 1, \dots, k$, temos

$$\begin{aligned} \frac{\partial^2 \ell(\beta, \gamma)}{\partial \beta_i \partial \beta_p} &= \frac{\partial}{\partial \beta_p} \left[\sum_{t=1}^n \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t} \frac{d\mu_t}{d\eta_t} \frac{\partial \eta_t}{\partial \beta_i} \right] \\ &= \sum_{t=1}^n \frac{\partial}{\partial \beta_p} \left[\frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t} \frac{d\mu_t}{d\eta_t} \right] \frac{\partial \eta_t}{\partial \beta_i} \\ &= \sum_{t=1}^n \frac{\partial}{\partial \mu_t} \left(\frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t} \frac{d\mu_t}{d\eta_t} \right) \frac{d\mu_t}{d\eta_t} \frac{\partial \eta_t}{\partial \beta_p} \frac{\partial \eta_t}{\partial \beta_i} + \sum_{t=1}^n \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t} \frac{d\mu_t}{d\eta_t} \frac{\partial^2 \eta_t}{\partial \beta_i \partial \beta_p} \\ &= \sum_{t=1}^n \left(\frac{\partial^2 \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t^2} \frac{d\mu_t}{d\eta_t} + \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t} \frac{\partial}{\partial \mu_t} \frac{d\mu_t}{d\eta_t} \right) \frac{d\mu_t}{d\eta_t} \frac{\partial \eta_t}{\partial \beta_i} \frac{\partial \eta_t}{\partial \beta_p} \\ &\quad + \sum_{t=1}^n \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t} \frac{d\mu_t}{d\eta_t} \frac{\partial^2 \eta_t}{\partial \beta_i \partial \beta_p}. \end{aligned}$$

De (2.5) e pela Proposição 2.1 (vii) temos que $\mathbb{E}(\partial \ell_t(\mu_t, \sigma_t^2)/\partial \mu_t) = 0$. Logo,

$$\mathbb{E} \left(\frac{\partial^2 \ell(\beta, \gamma)}{\partial \beta_i \partial \beta_p} \right) = \sum_{t=1}^n \mathbb{E} \left(\frac{\partial^2 \ell_t(\mu_t, \sigma_t^2)}{\partial \mu_t^2} \right) \left(\frac{d\mu_t}{d\eta_t} \right)^2 \frac{\partial \eta_t}{\partial \beta_i} \frac{\partial \eta_t}{\partial \beta_p}. \quad (2.10)$$

De (2.5) temos que

$$\frac{\partial^2 \ell_t(\mu_t, \sigma^2)}{\partial \mu_t^2} = -\frac{1}{2\sigma_t^2} d''(y_t; \mu_t).$$

Usando a Proposição 2.1 (v), ou seja,

$$\frac{1}{2} \mathbb{E}[d''(y_t, \mu_t)] = \frac{3\sigma_t^2}{\mu_t(1 - \mu_t)} + \frac{1}{\mu_t^3(1 - \mu_t)^3},$$

podemos reescrever (2.10) como

$$\mathbb{E} \left(\frac{\partial^2 \ell(\beta, \gamma)}{\partial \beta_i \partial \beta_p} \right) = - \sum_{t=1}^n \frac{1}{\sigma_t^2} w_t \frac{\partial \eta_t}{\partial \beta_i} \frac{\partial \eta_t}{\partial \beta_p},$$

em que

$$w_t = \left(\frac{3\sigma_t^2}{\mu_t(1-\mu_t)} + \frac{1}{\mu_t^3(1-\mu_t)^3} \right) \frac{1}{[g'(\mu_t)]^2}. \quad (2.11)$$

Portanto, a informação de Fisher para β pode ser escrita matricialmente da seguinte forma:

$$K_{\beta\beta} = -\mathbb{E} \left(\frac{\partial^2 \ell(\beta, \gamma)}{\partial \beta_i \partial \beta_p} \right) = \tilde{\mathcal{X}}^\top S W \tilde{\mathcal{X}},$$

em que $W = \text{diag}\{w_1, \dots, w_n\}$ e a matriz S está definida em (2.7).

Considerando agora as derivadas de segunda ordem do logaritmo da função de verossimilhança (2.4) com respeito a β_i e γ_j , $i = 1, \dots, k$ e $j = 1, \dots, q$, temos que

$$\begin{aligned} \frac{\partial^2 \ell(\beta, \gamma)}{\partial \beta_i \partial \gamma_j} &= \sum_{t=1}^n \frac{\partial}{\partial \beta_i} \left[\frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \gamma_j} \right] \\ &= \sum_{t=1}^n \frac{\partial}{\partial \beta_i} \left[\left(-\frac{1}{2\sigma_t^2} + \frac{d(y_t; \mu_t)}{2(\sigma_t^2)^2} \right) \frac{1}{h'(\sigma_t^2)} \frac{\partial \zeta_t}{\partial \gamma_j} \right] \\ &= \sum_{t=1}^n \left(\frac{d'(y_t; \mu_t)}{2(\sigma_t^2)^2} \right) \frac{1}{g'(\mu_t)} \frac{1}{h'(\sigma_t^2)} \frac{\partial \eta_t}{\partial \beta_i} \frac{\partial \zeta_t}{\partial \gamma_j}. \end{aligned}$$

De acordo com a Proposição 2.1 (vii) temos que $\mathbb{E}[d'(y_t; \mu_t)] = 0$. Assim,

$$\mathbb{E} \left(\frac{\partial^2 \ell(\beta, \gamma)}{\partial \beta_i \partial \gamma_j} \right) = 0,$$

e, portanto, $K_{\beta\gamma} = K_{\gamma\beta}^\top = 0$.

Por último, a derivada de segunda ordem do logaritmo da função de verossimilhança (2.4) com relação a γ_j e γ_l , para $j, l = 1, \dots, q$, é dada por

$$\begin{aligned} \frac{\partial^2 \ell(\beta, \gamma)}{\partial \gamma_j \partial \gamma_l} &= \frac{\partial}{\partial \gamma_l} \left[\sum_{t=1}^n \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \sigma_t^2} \frac{d\sigma_t^2}{d\zeta_t} \frac{\partial \zeta_t}{\partial \gamma_j} \right] \\ &= \sum_{t=1}^n \frac{\partial}{\partial \gamma_l} \left[\frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \sigma_t^2} \frac{d\sigma_t^2}{d\zeta_t} \right] \frac{\partial \zeta_t}{\partial \gamma_j} \\ &= \sum_{t=1}^n \frac{\partial}{\partial \sigma_t^2} \left(\frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \sigma_t^2} \frac{d\sigma_t^2}{d\zeta_t} \right) \frac{d\sigma_t^2}{d\zeta_t} \frac{\partial \zeta_t}{\partial \gamma_l} \frac{\partial \zeta_t}{\partial \gamma_j} + \sum_{t=1}^n \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \sigma_t^2} \frac{d\sigma_t^2}{d\zeta_t} \frac{\partial^2 \zeta_t}{\partial \gamma_j \partial \gamma_l} \\ &= \sum_{t=1}^n \left(\frac{\partial^2 \ell_t(\mu_t, \sigma_t^2)}{\partial (\sigma_t^2)^2} \frac{d\sigma_t^2}{d\zeta_t} + \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \sigma_t^2} \frac{\partial}{\partial \sigma_t^2} \frac{d\sigma_t^2}{d\zeta_t} \right) \frac{d\sigma_t^2}{d\zeta_t} \frac{\partial \zeta_t}{\partial \gamma_j} \frac{\partial \zeta_t}{\partial \gamma_l} \\ &\quad + \sum_{t=1}^n \frac{\partial \ell_t(\mu_t, \sigma_t^2)}{\partial \sigma_t^2} \frac{d\sigma_t^2}{d\zeta_t} \frac{\partial^2 \zeta_t}{\partial \gamma_j \partial \gamma_l}. \end{aligned}$$

Usando a Proposição 2.1 (i) e a Equação (2.8) temos que $\mathbb{E}[\partial \ell_t(\mu_t, \sigma_t^2)/\partial \sigma_t^2] = 0$. Então,

$$\mathbb{E} \left(\frac{\partial^2 \ell(\beta, \gamma)}{\partial \gamma_j \partial \gamma_l} \right) = \sum_{t=1}^n \mathbb{E} \left(\frac{\partial^2 \ell_t(\mu_t, \sigma_t^2)}{\partial (\sigma_t^2)^2} \right) \left(\frac{d\sigma_t^2}{d\zeta_t} \right)^2 \frac{\partial \zeta_t}{\partial \gamma_j} \frac{\partial \zeta_t}{\partial \gamma_l}.$$

Adicionalmente, de (2.8) obtemos

$$\frac{\partial^2 \ell_t(\mu_t, \sigma_t^2)}{\partial (\sigma_t^2)^2} = \frac{1}{2(\sigma_t^2)^2} - \frac{d(y_t; \mu_t)}{(\sigma_t^2)^3}.$$

Logo,

$$\mathbb{E} \left(\frac{\partial^2 \ell(\beta, \gamma)}{\partial \gamma_j \partial \gamma_l} \right) = - \sum_{t=1}^n d_t \frac{\partial \zeta_t}{\partial \gamma_j} \frac{\partial \zeta_t}{\partial \gamma_l},$$

em que

$$d_t = \frac{1}{2(\sigma_t^2)^2} \frac{1}{[h'(\sigma_t^2)]^2}.$$

Assim, a informação de Fisher para γ é dada matricialmente por

$$K_{\gamma\gamma} = -\mathbb{E} \left(\frac{\partial^2 \ell(\beta, \gamma)}{\partial \gamma_j \partial \gamma_l} \right) = \tilde{\mathbf{Z}}^\top D \tilde{\mathbf{Z}},$$

em que $D = \text{diag}(d_1, \dots, d_n)$.

Portanto, a matriz de informação de Fisher para o vetor de parâmetros $\theta = (\beta^\top, \gamma^\top)^\top$ é dada por

$$K = K(\beta, \gamma) = \begin{pmatrix} K_{\beta\beta} & 0 \\ 0 & K_{\gamma\gamma} \end{pmatrix}, \quad (2.12)$$

em que $K_{\beta\beta} = \tilde{\mathbf{X}}^\top S W \tilde{\mathbf{X}}$ e $K_{\gamma\gamma} = \tilde{\mathbf{Z}}^\top D \tilde{\mathbf{Z}}$. Como K é uma matriz bloco diagonal, os vetores de parâmetros β e γ são globalmente ortogonais (COX; REID, 1987), de modo que seus estimadores de máxima verossimilhança $\hat{\beta}$ e $\hat{\gamma}$, respectivamente, são assintoticamente independentes (ESPINHEIRA; SILVA, 2019).

Para grandes amostras e sob condições de regularidade (ver Apêndice B), a distribuição aproximada dos estimadores de máxima verossimilhança é dada por

$$\begin{pmatrix} \hat{\beta} \\ \hat{\gamma} \end{pmatrix} \sim N_{k+q} \left(\begin{pmatrix} \beta \\ \gamma \end{pmatrix}, K^{-1} \right), \quad (2.13)$$

em que K^{-1} é a inversa da matriz de informação de Fisher definida por

$$K^{-1} = K(\beta, \gamma)^{-1} = \begin{pmatrix} K^{\beta\beta} & 0 \\ 0 & K^{\gamma\gamma} \end{pmatrix}, \quad (2.14)$$

em que $K^{\beta\beta} = (\tilde{\mathbf{X}}^\top S W \tilde{\mathbf{X}})^{-1}$ e $K^{\gamma\gamma} = (\tilde{\mathbf{Z}}^\top D \tilde{\mathbf{Z}})^{-1}$.

3 ESTIMAÇÃO PONTUAL E INTERVALAR

3.1 INTRODUÇÃO

A estimação dos parâmetros que indexam a distribuição simplex é tipicamente feita por máxima verossimilhança. Estes estimadores podem ser viesados quando o tamanho da amostra é pequeno ou até mesmo moderado. Desta forma, é desejável a obtenção de estimadores com viés reduzido em amostras finitas. Assim, o método bootstrap (EFRON, 1979) surge como uma opção para o melhoramento inferencial, pois através dele é possível construir estimadores pontuais corrigidos com menor viés.

Na inferência estatística é de fundamental importância associar confiabilidade às estimativas pontuais do modelo, e uma das formas de fazer isto é através da construção das estimativas intervalares dos parâmetros. Na literatura, existem várias abordagens para a construção de intervalos de confiança. A mais utilizada é o intervalo de confiança assintótico, que supõe normalidade assintótica dos estimadores de máxima verossimilhança. Esse intervalo requer grandes amostras para que a cobertura exata esteja próxima da nominal. Uma alternativa para a construção de intervalos de confiança adequados em pequenas amostras é através do método bootstrap. Existem diversas técnicas para o cálculo de intervalos de confiança bootstrap. Neste trabalho, faremos uso dos intervalos de confiança bootstrap percentil e bootstrap- t (EFRON; TIBSHIRANI, 1994).

O objetivo deste capítulo é corrigir o viés via esquema bootstrap dos estimadores de máxima verossimilhança do modelo de regressão simplex não linear em pequenas amostras. Além disso, propor melhorias para a estimação intervalar dos parâmetros do modelo através do método bootstrap. Posteriormente, avaliar numericamente via simulações de Monte Carlo as correções propostas e compará-las com as versões não corrigidas. Por último, aplicar os melhoramentos inferenciais propostos ao conjunto de dados do departamento de Química da Universidade Nacional da Colômbia (SALAZAR, 2005), composto por 28 observações. A aplicação consiste no processo de craqueamento catalítico fluido (SALAZAR, 2005).

3.2 ESTIMAÇÃO PONTUAL DOS PARÂMETROS DO MODELO

Os estimadores de máxima verossimilhança de β e γ são obtidos como solução do sistema $U_{\beta}(\beta, \gamma) = 0$ e $U_{\gamma}(\beta, \gamma) = 0$. Porém, tais estimadores não possuem forma fechada, sendo

necessário a utilização de métodos numéricos para a estimação dos parâmetros.

Um método bastante utilizado para a obtenção do estimador de máxima verossimilhança $\hat{\theta}$ de θ é o método iterativo de Newton-Raphson, que se baseia em expandir em série de Taylor a função escore U , em torno de um valor inicial arbitrário. Dessa forma, é possível chegar ao seguinte processo iterativo

$$\theta^{(k+1)} = \theta^{(k)} + [-(U')^{-1}]^{(k)} U^{(k)}, \quad k = 0, 1, \dots,$$

em que $-(U')$ é a matriz de informação observada de θ . Porém, nem sempre a matriz $-(U')$ é positiva definida quando θ estiver longe de seu valor ótimo, então, uma solução é substituir $-(U')$ pelo seu correspondente valor esperado, ou seja, pela matriz de informação de Fisher K , definida em (2.12). Assim, temos o processo iterativo Escore de Fisher, dado por

$$\theta^{(k+1)} = \theta^{(k)} + [K^{-1}]^{(k)} U^{(k)}, \quad k = 0, 1, \dots$$

Esse processo é repetido até que a diferença absoluta entre $\theta^{(k+1)}$ e $\theta^{(k)}$ seja menor que uma tolerância especificada.

É necessário atribuir um valor para $\theta = (\beta^\top, \gamma^\top)^\top$ com o objetivo de inicializar o processo iterativo. O chute inicial proposto em Espinheira e Silva (2019) para β e γ no modelo de regressão simplex não linear é baseado em mínimos quadrados não lineares e uma aproximação não linear.

Suponha que $k_1 = k$ e $q_1 = q$. A expansão de Taylor de primeira ordem de $f(x_t; \beta)$ em torno de $\beta^{(0)}$ é tal que

$$f(x_t^\top; \beta) \approx f(x_t^\top; \beta^{(0)}) + \sum_{i=1}^k \left[\frac{\partial f(x_t^\top; \beta)}{\partial \beta_i} \right]_{\beta=\beta^{(0)}} (\beta_i - \beta_i^{(0)}), \quad t = 1, \dots, n,$$

em que $\beta^{(0)} = (\beta_1^{(0)}, \dots, \beta_k^{(0)})^\top$ é um valor inicial. Assim, os autores propõem considerar

$$f(x_t^\top; \beta) = f(x_t^\top; \beta^{(0)}) + \sum_{i=1}^k \tilde{x}_{ti}^{(0)} (\beta_i - \beta_i^{(0)}). \quad (3.1)$$

No caso do modelo normal não linear temos que $y_t = f_1(x_t^\top; \beta) + \xi_t$, tal que $\mathbb{E}(y_t | x_t^\top) = \mu_t$ e na Equação (3.1) é considerada a aproximação $y_t \approx \mu_t = f_1(x_t^\top; \beta)$, ou seja, $y_t \approx f_1(x_t^\top; \beta)$. Aqui, vamos considerar $g(y_t) \approx g(\mu_t) = f_1(x_t^\top; \beta)$. Seja $\theta_i^{(0)} = (\beta_i - \beta_i^{(0)})$. Sem perda de generalidade vamos considerar então $g(y_t) - f(x_t^\top; \beta^{(0)}) = \sum_{i=1}^k \tilde{x}_{ti}^{(0)} \theta_i^{(0)}$, $t = 1, \dots, n$, que pode ser visto como um modelo linear para o qual o estimador de mínimos quadrados de $\theta^{(0)}$ é dado por $\hat{\theta}^{(0)} = (\tilde{\mathcal{X}}^{(0)\top} \tilde{\mathcal{X}}^{(0)})^{-1} \tilde{\mathcal{X}}^{(0)\top} [g(y) - f(x^\top; \beta^{(0)})]$, com $\tilde{\mathcal{X}}^{(0)} = [\partial \eta / \partial \beta]_{\beta=\beta^{(0)}}$ e

$\hat{\theta}_i^{(0)} = (\hat{\beta}_i^{(1)} - \beta_i^{(0)})$. Consequentemente, $\hat{\beta}_i^{(1)} = \hat{\theta}_i^{(0)} + \beta_i^{(0)}$. A proposta de Espinheira e Silva (2019) é usar o seguinte chute inicial não linear para β :

$$\beta_{NL}^{(0)} = (\tilde{\mathcal{X}}^{(0)\top} \tilde{\mathcal{X}}^{(0)})^{-1} \tilde{\mathcal{X}}^{(0)\top} f_1(X; \beta_L^{(0)}),$$

em que $f_1(x_t^\top; \beta_L^{(0)}) = \eta_t$, $t = 1, \dots, n$, avaliado em $\beta^{(0)} = \beta_L^{(0)} = (X^\top X)^{-1} X^\top g(y)$. Como se pode notar, o procedimento proposto para selecionar os pontos de partida envolve dois passos. O primeiro é obter $\beta_L^{(0)}$ e o segundo calcular o chute inicial não linear $\beta_{NL}^{(0)}$.

Para o submodelo da dispersão os autores consideram $h(\sigma_t^2) = f(z_t^\top; \gamma)$. Assim,

$$\gamma_{NL}^{(0)} = (\tilde{\mathcal{Z}}^{(0)\top} \tilde{\mathcal{Z}}^{(0)})^{-1} \tilde{\mathcal{Z}}^{(0)\top} [h(\sigma_{NL}^2) - f_2(Z; \gamma_L^{(0)})],$$

em que $f_2(z_t^\top; \gamma_L^{(0)}) = \zeta_t$ avaliado em $\gamma^{(0)} = \gamma_L^{(0)} = (Z^\top Z)^{-1} Z^\top h(\sigma_L^2)$. Adicionalmente, temos que $\tilde{\mathcal{Z}}^{(0)} = [\partial \zeta / \partial \gamma]_{\gamma=\gamma^{(0)}}$, $\sigma_L^{2(0)} = (\sigma_{L1}^{2(0)}, \sigma_{L2}^{2(0)}, \dots, \sigma_{Ln}^{2(0)})^\top$ em que $\sigma_{Lt}^2 = d(y_t; \check{\mu}_{Lt})$ com $\check{\mu}_{Lt} = g^{-1}(\hat{\eta}_{Lt}) = g^{-1}(x_t^\top \beta_L^{(0)})$. Finalmente,

$$\sigma_{NLt}^{2(0)} = d(y_t; \check{\mu}_{NLt}),$$

em que $\check{\mu}_{NLt} = g^{-1}(\hat{\eta}_{NLt}) = g^{-1}[f(x_t^\top; \beta_{NL}^{(0)})]$.

De acordo com (ESPINHEIRA; SILVA, 2019), normalmente, em um ou nos dois submodelos há mais parâmetros que covariáveis nos preditores não lineares, isto é, $k_1 < k$ e/ou $q_1 < q$. Nesses casos, é necessário fornecer valores numéricos para os parâmetros para obter um preditor que não envolva parâmetros desconhecidos. Este passo é relevante, uma vez que estas estimativas numéricas iniciais devem respeitar as características das covariáveis e sua relação com as funções matemáticas que descrevem os preditores não lineares. Após este passo é construída uma matriz X baseada em um preditor linear e é calculado $(X^\top X)^{-1} X^\top g(y)$. Com base neste procedimento de dois passos, é possível obter $\beta_L^{(0)}$ quando o número de parâmetros exceder o número de covariáveis.

3.3 MÉTODO BOOTSTRAP

Conforme mencionado anteriormente, os estimadores de máxima verossimilhança podem ser viesados quando o tamanho da amostra é pequeno ou até mesmo moderado. Assim, é desejável a obtenção de estimadores com viés reduzido em amostras finitas.

O viés mede a distância média entre o estimador e o verdadeiro valor do parâmetro. Frequentemente, o viés é utilizado para avaliar a qualidade de um estimador. Em grandes

amostras, esse viés não é tipicamente problemático, pois em geral é de ordem $O(n^{-1})$ enquanto o respectivo erro-padrão é de ordem $O(n^{-1/2})$, em que n é o tamanho da amostra (OSPINA; CRIBARI-NETO; VASCONCELLOS, 2006). Todavia, em pequenas amostras, a existência do viés pode ser problemática (EFRON; TIBSHIRANI, 1994; PAWITAN, 2001; OSPINA; CRIBARI-NETO; VASCONCELLOS, 2006). Este problema pode ser contornado através do método bootstrap.

O bootstrap é um método computacional de inferência estatística introduzido em 1979 por Efron (EFRON, 1979), capaz de responder a questões reais sem a necessidade de exaustivos, complicados e, muitas vezes, inviáveis cálculos analíticos. O bootstrap é conhecido como uma das possíveis técnicas de reamostragem, que é o nome que se dá, a um conjunto de técnicas que se baseiam em calcular estimativas a partir de repetidas amostragens dentro da mesma amostra. Este procedimento é de grande utilidade para casos em que as técnicas tradicionais não podem ser aplicadas ou são muito complicadas. Este método é uma importante ferramenta para contornar problemas de inferências em tamanho de amostra pequeno ou moderado. O mesmo apresenta diversas aplicações, por exemplo, a correção de viés de estimadores, a construção de intervalos de confiança, a realização de testes de hipóteses, entre outras.

Os dois tipos de métodos bootstrap utilizados são o paramétrico, quando se supõe conhecido o modelo de probabilidade associado aos dados, e o não-paramétrico, quando não é possível estabelecer essa suposição. Dessa forma, o método bootstrap pode ser implementado tanto de forma paramétrica quanto não-paramétrica. O que difere os dois métodos é a forma de obtenção das amostras artificiais. No caso paramétrico, desde que a distribuição dos dados seja conhecida, a amostra será composta realizando-se a amostragem diretamente dessa distribuição e os parâmetros desconhecidos serão substituídos por estimativas pontuais. Por outro lado, no caso não-paramétrico, a amostra será composta por extrações dos elementos da amostra original com reposição. Aqui, vamos utilizar o método bootstrap paramétrico, pois no nosso caso assumimos um modelo paramétrico conhecido para os dados observados. Os passos para a realização deste método estão descritos no Algoritmo 1.

Algoritmo 1: Método paramétrico

- 1: Suponha que $y = (y_1, \dots, y_n)^\top$ é uma amostra aleatória em que cada y_t , $t = 1, \dots, n$, segue distribuição F indexada pelo vetor de parâmetros θ ;
 - 2: A partir da amostra original, obter a estimativa $\hat{\theta}$ de θ ;
 - 3: Gerar B amostras bootstrap de tamanho n ($y_1^*, \dots, y_n^*$) de $F(\hat{\theta})$;
 - 4: Para cada amostra bootstrap y_b^* calcular $\hat{\theta}^{*b}$;
 - 5: Repita os passos 3 e 4 um número muito grande B de vezes, obtendo assim $\hat{\theta}^{*1}, \dots, \hat{\theta}^{*B}$;
 - 6: Use as estimativas $\hat{\theta}^{*b}$, com $b = 1, \dots, B$, para calcular as quantidades desejadas (média, variância, intervalo de confiança, etc.).
-

3.4 CORREÇÃO DE VIÉS DOS ESTIMADORES PONTUAIS

Através do método bootstrap podemos estimar o viés de um estimador pontual. Uma vez obtida uma estimativa do viés do estimador podemos construir estimadores pontuais corrigidos pelo viés. Usando as etapas do método bootstrap apresentadas no Algoritmo 1, uma estimativa bootstrap do viés pode ser obtida por

$$\hat{B}_{boot}(\hat{\theta}) = \bar{\theta}^* - \hat{\theta},$$

em que $\bar{\theta}^* = \frac{1}{B} \sum_{b=1}^B \hat{\theta}^{*b}$, ou seja, é possível aproximar o valor esperado a partir da média aritmética das estimativas bootstrap de θ . Assim, podemos obter um estimador corrigido até segunda ordem por bootstrap (EFRON, 1979; DAVISON; HINKLEY, 1997):

$$\bar{\theta} = \hat{\theta} - \hat{B}_{boot}(\hat{\theta}) = \hat{\theta} - (\bar{\theta}^* - \hat{\theta}) = 2\hat{\theta} - \bar{\theta}^*.$$

Este estimador possui as mesmas propriedades assintóticas que o EMV usual, mas apresenta menor viés em pequenas amostras (EFRON; TIBSHIRANI, 1994). Uma discussão detalhada sobre a correção de viés de segunda ordem por bootstrap e sua relação com a correção analítica pode ser encontrada em Ferrari e Cribari-Neto (1998).

3.5 ESTIMAÇÃO INTERVALAR DOS PARÂMETROS DO MODELO

Sejam y a variável aleatória de interesse, $F(y; \theta)$ a distribuição de y , $\mathcal{N} = \{F(y; \theta), \theta \in \Theta \subset \mathbb{R}^{(k+q)}\}$ a família de distribuição a qual a distribuição de y pertence e Θ o espaço paramétrico. Considere θ o vetor de parâmetros da distribuição da variável aleatória y e seja $\hat{\theta}$ o EMV de θ .

A estimação intervalar com confiança baseia-se na determinação de um conjunto construído

a partir de uma probabilidade pré-estabelecida deste conjunto conter o verdadeiro valor do parâmetro. Uma estimativa intervalar de confiança para o parâmetro θ é um intervalo $[l_1, l_2]$, em que $l_1 < l_2$ tal que a probabilidade do intervalo aleatório conter o parâmetro de interesse θ é igual ao valor pré-fixado $1 - \alpha$, com $0 < \alpha < 1$, ou seja,

$$\mathbb{P}[l_1 \leq \theta \leq l_2] = 1 - \alpha,$$

em que $1 - \alpha$ é o nível de confiança do intervalo, isto é, $1 - \alpha$ é a probabilidade de cobertura e l_1 e l_2 são os limites inferior e superior do intervalo de confiança, respectivamente. Intervalos de confiança com probabilidade de cobertura igual ao nível de confiança, como apresentado acima, são ditos exatos. Porém, as vezes não é possível construir intervalos exatos e assim intervalos de confiança aproximados são utilizados, ou seja,

$$\mathbb{P}[l_1 \leq \theta \leq l_2] \approx 1 - \alpha.$$

Assim, o intervalo é dito aproximado, exatamente porque os limites l_1 e l_2 são aproximados. Vale enfatizar, que o intervalo é a quantidade aleatória, não o parâmetro.

Existem várias abordagens para a construção de intervalos de confiança aproximados. A mais utilizada é o intervalo de confiança assintótico, que supõe normalidade assintótica dos EMVs. De acordo com (2.13) para o modelo simplex a distribuição assintótica de $\hat{\theta} = (\hat{\beta}^\top, \hat{\gamma}^\top)^\top$ é aproximadamente normal com média igual a $\theta = (\beta^\top, \gamma^\top)^\top$ e matriz de variâncias e covariâncias dada por (2.14). Mais especificamente, temos que $K^{\beta\beta} = (\tilde{\mathcal{X}}^\top S W \tilde{\mathcal{X}})^{-1}$ avaliado em $\hat{\theta}$ é a matriz $k \times k$ de variâncias e covariâncias de $\hat{\beta}$ e $K^{\gamma\gamma} = (\tilde{\mathcal{Z}}^\top D \tilde{\mathcal{Z}})^{-1}$ avaliado em $\hat{\theta}$ é a matriz $q \times q$ de variâncias e covariâncias de $\hat{\gamma}$.

Considere β_i e γ_j , com $i = 1, \dots, k$ e $j = 1, \dots, q$, o i -ésimo e o j -ésimo componentes dos vetores β e γ , respectivamente. Temos que $K_i^{\beta\beta}$ e $K_j^{\gamma\gamma}$ são os i -ésimo e j -ésimo componentes da diagonal principal das matrizes $K^{\beta\beta}(\hat{\theta})$ e $K^{\gamma\gamma}(\hat{\theta})$, respectivamente. Assim, segue que

$$\left(\hat{\beta}_i - z_{1-\frac{\alpha}{2}} (K_i^{\beta\beta})^{1/2}; \hat{\beta}_i + z_{1-\frac{\alpha}{2}} (K_i^{\beta\beta})^{1/2} \right)$$

e

$$\left(\hat{\gamma}_j - z_{1-\frac{\alpha}{2}} (K_j^{\gamma\gamma})^{1/2}; \hat{\gamma}_j + z_{1-\frac{\alpha}{2}} (K_j^{\gamma\gamma})^{1/2} \right)$$

são intervalos com confiança aproximadamente igual a $1 - \alpha$ para β_i e γ_j , respectivamente, em que $z_{1-\frac{\alpha}{2}}$ é o quantil $1 - \frac{\alpha}{2}$ da distribuição normal padrão.

Em inferência por máxima verossimilhança, esses intervalos podem requerer grandes amostras para que as coberturas exatas estejam próximas às nominais. Em pequenas amostras, eles podem apresentar grandes erros de cobertura (EFRON, 1979; DAVISON; HINKLEY, 1997).

Uma alternativa para a construção de intervalos de confiança adequados em pequenas amostras, livre de complexidades analíticas, é o método bootstrap, pois é possível construir intervalos de confiança que apresentem níveis de cobertura próximos da verdadeira probabilidade de cobertura. Existem diversas técnicas para o cálculo de intervalos de confiança via esquema bootstrap. Neste trabalho, faremos uso dos intervalos de confiança bootstrap percentil e bootstrap- t .

3.5.1 Intervalo de Confiança Bootstrap Percentil

Através do método bootstrap paramétrico ou não paramétrico, são geradas amostras bootstrap $(y^{*1}, y^{*2}, \dots, y^{*B})$ e a partir delas obtêm-se as réplicas bootstrap de $\hat{\theta}$, $\hat{\theta}^{*1}, \hat{\theta}^{*2}, \dots, \hat{\theta}^{*B}$. Seja $\hat{G}$ a função de distribuição empírica de $\hat{\theta}$ obtida através das B réplicas bootstrap. Podemos construir o intervalo de confiança bootstrap percentil, com nível aproximado de cobertura $1 - \alpha$, definindo os quantis $\alpha/2$ e $1 - \alpha/2$ de $\hat{G}$ da seguinte forma:

$$(\hat{G}^{-1}(\alpha/2), \hat{G}^{-1}(1 - \alpha/2)).$$

Definindo $\hat{G}^{-1}(\alpha/2) = \hat{\theta}^{*(\alpha/2)}$ e $\hat{G}^{-1}(1 - \alpha/2) = \hat{\theta}^{*(1-\alpha/2)}$ podemos escrever o intervalo percentil como

$$(\hat{\theta}^{*(\alpha/2)}, \hat{\theta}^{*(1-\alpha/2)}). \quad (3.2)$$

Desta forma, os intervalos percentil para os parâmetros do modelo de regressão simplex não linear são dados pelas seguintes expressões

$$(\hat{\beta}_i^{*(\alpha/2)}, \hat{\beta}_i^{*(1-\alpha/2)}) \quad \text{e} \quad (\hat{\gamma}_j^{*(\alpha/2)}, \hat{\gamma}_j^{*(1-\alpha/2)}).$$

O intervalo percentil pode ser obtido conforme o Algoritmo 2.

O intervalo percentil dado em (3.2) não necessariamente é simétrico em relação ao valor de $\hat{\theta}$. A sua construção garante a não inclusão de valores impróprios para o parâmetro de interesse no intervalo de confiança. Além disso, outra vantagem do método percentil é que através de uma transformação podemos normalizar perfeitamente a distribuição de $\hat{\theta}$ (EFRON; TIBSHIRANI, 1994).

Algoritmo 2: Intervalo de confiança bootstrap percentil

- 1: Gera-se B amostras bootstrap $(y^{*1}, y^{*2}, \dots, y^{*B})$;
 - 2: Calcula-se as respectivas estimativas $\hat{\theta}^{*b} = s(y^{*b})$, $b = 1, \dots, B$; $\hat{\theta}^{*b} = s(y^{*b})$ é o valor estimado de θ para a amostra bootstrap y^{*b} ;
 - 3: Ordena-se as B réplicas de $\hat{\theta}$;
 - 4: Considera-se como limites inferior e superior do intervalo percentil as réplicas de $\hat{\theta}$ de ordem $B \times (\alpha/2)$ e $B \times (1 - \alpha/2)$, respectivamente, assumindo que $B \times (\alpha/2)$ e $B \times (1 - \alpha/2)$ são inteiros e $0 < \alpha < 1/2$;
 - 4.1: Se $B \times (\alpha/2)$ e $B \times (1 - \alpha/2)$ não forem inteiros, podemos utilizar o seguinte procedimento:
 - 4.1.1: Assumindo que $0 < \alpha < 1/2$, seja $k = [(B + 1) \times \alpha/2]$ o maior inteiro menor ou igual ao número $(B + 1) \times \alpha/2$; então, definimos os limites inferior e superior do intervalo percentil pelos k -ésimo e $(B + 1 - k)$ -ésimo elementos ordenados das B réplicas bootstrap, respectivamente.
-

Suponha que $\hat{\theta}$ não tem distribuição normal em amostras finitas e que h é uma função que normaliza a distribuição de $\hat{\theta}$, isto é,

$$\hat{\phi} \sim \mathcal{N}(\phi, c^2),$$

com $\hat{\phi} = h(\hat{\theta})$, em que c é algum desvio-padrão constante. Desta forma, podemos construir um intervalo de confiança exato com probabilidade de cobertura $1 - \alpha$ para o parâmetro ϕ da forma

$$\left(\hat{\phi} - z_{1-\frac{\alpha}{2}} c, \hat{\phi} + z_{1-\frac{\alpha}{2}} c \right),$$

em que $z_{1-\frac{\alpha}{2}}$ é o percentil $(1 - \alpha/2)$ da distribuição normal padrão. Assim, o intervalo percentil de $\hat{\theta}$ pode ser dado como

$$\left(h^{-1}(\hat{\phi} - z_{1-\frac{\alpha}{2}} c), h^{-1}(\hat{\phi} + z_{1-\frac{\alpha}{2}} c) \right),$$

sem a necessidade de conhecer a função normalizadora $h(\cdot)$. Para mais detalhes sobre o intervalo de confiança bootstrap percentil, ver Efron e Tibshirani (1994).

3.5.2 Intervalo de Confiança Bootstrap- t

O intervalo bootstrap- t , também conhecido como método pivotal, é uma generalização do método t -Student, e é geralmente aplicado em estatísticas de posição/localização como a média amostral, mediana ou percentil amostral. Este intervalo é obtido através da estimativa da distribuição T diretamente da amostra aleatória observada $y = (y_1, y_2, \dots, y_n)^\top$, em que

T é dada por

$$T = \frac{\hat{\theta} - \theta}{\widehat{ep}(\hat{\theta})},$$

sendo T conhecida como estatística t quando há suposição de normalidade nos dados e $\widehat{ep}(\hat{\theta})$ sendo o erro-padrão de $\hat{\theta}$. Os passos para a construção do intervalo de confiança bootstrap- t estão descritos no Algoritmo 3.

Algoritmo 3: Intervalo de confiança bootstrap- t

- 1: Gera-se B amostras bootstrap $(y^{*1}, y^{*2}, \dots, y^{*B})$;
- 2: Para cada amostra bootstrap, calcula-se

$$T^{*b} = \frac{\hat{\theta}^{*b} - \hat{\theta}}{\widehat{ep}^{*b}},$$

- com $b = 1, 2, \dots, B$, em que $\hat{\theta} = s(y)$ é o valor estimado de θ a partir da amostra original y , $\hat{\theta}^{*b} = s(y^{*b})$ é o valor estimado de θ para a amostra bootstrap y^{*b} e $\widehat{ep}^{*b}$ é o erro-padrão de $\hat{\theta}^{*b}$ para a amostra bootstrap y^{*b} . Nota-se que $ep = \kappa(\theta)$, κ função conhecida e $\widehat{ep}^{*b} = \kappa(\hat{\theta}^{*b})$;
- 3: Os percentis $\alpha/2$ e $1 - \alpha/2$ de T^{*b} são estimados pelos valores $\hat{t}^{(\alpha/2)}$ e $\hat{t}^{(1-\alpha/2)}$, respectivamente, tais que

$$\frac{\#\{T^{*b} \leq \hat{t}^{(\alpha/2)}\}}{B} = \frac{\alpha}{2} \quad \text{e} \quad \frac{\#\{T^{*b} \leq \hat{t}^{(1-\alpha/2)}\}}{B} = 1 - \frac{\alpha}{2}.$$

Desta forma, o intervalo de confiança bootstrap- t é dado por

$$(\hat{\theta} - \hat{t}^{(1-\alpha/2)}\widehat{ep}, \hat{\theta} - \hat{t}^{(\alpha/2)}\widehat{ep}),$$

em que $\widehat{ep} \equiv \widehat{ep}(\hat{\theta})$.

As quantidades $\hat{t}^{(\alpha/2)}$ e $\hat{t}^{(1-\alpha/2)}$ podem ser obtidas da seguinte forma:

1. Ordene as B réplicas bootstrap T^{*b} ;
2. As réplicas de ordem $B \times (\alpha/2)$ e $B \times (1 - \alpha/2)$ são as quantidades $\hat{t}^{(\alpha/2)}$ e $\hat{t}^{(1-\alpha/2)}$, respectivamente, assumindo que $B \times (\alpha/2)$ e $B \times (1 - \alpha/2)$ são inteiros;

Se $B \times (\alpha/2)$ e $B \times (1 - \alpha/2)$ não forem inteiros, podemos utilizar o seguinte procedimento:

Assumindo que $0 < \alpha < 1/2$, seja $k = [(B + 1) \times \alpha/2]$ o maior inteiro menor ou igual ao número $(B + 1) \times \alpha/2$; então, as quantidades $\hat{t}^{(\alpha/2)}$ e $\hat{t}^{(1-\alpha/2)}$ são dadas pelos k -ésimo e $(B + 1 - k)$ -ésimo elementos ordenados de T^{*b} , respectivamente.

Portanto, os intervalos bootstrap- t para os parâmetros do modelo de regressão simplex

não linear, são dados pelas seguintes expressões

$$(\hat{\beta}_i - \hat{t}^{(1-\alpha/2)} \widehat{ep}, \hat{\beta}_i - \hat{t}^{(\alpha/2)} \widehat{ep}) \quad \text{e} \quad (\hat{\gamma}_j - \hat{t}^{(1-\alpha/2)} \widehat{ep}, \hat{\gamma}_j - \hat{t}^{(\alpha/2)} \widehat{ep}).$$

Para mais informações sobre a construção do intervalo de confiança bootstrap- t , ver Efron e Tibshirani (1994). Vale ressaltar que a utilização do método bootstrap não implica que os demais métodos devem ser ignorados.

3.6 RESULTADOS NUMÉRICOS

Nesta seção apresentamos os resultados de simulações de Monte Carlo realizadas para avaliar os desempenhos dos estimadores de máxima verossimilhança do modelo de regressão simplex não linear e suas versões bootstrap em pequenas amostras.

Nas avaliações numéricas, consideramos o seguinte modelo de regressão simplex não linear:

$$g(\mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4} \quad \text{e} \quad h(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}, \quad t = 1, \dots, n, \quad (3.3)$$

em que $g(\cdot)$ e $h(\cdot)$ são as funções de ligação logito e logarítmica, respectivamente. O modelo escolhido é um dos mais simples da classe dos modelos não lineares.

As realizações das covariáveis foram geradas através da distribuição uniforme como segue: $x_{t2} \sim \mathcal{U}(0.5, 1.5)$, $x_{t3} \sim \mathcal{U}(0, 1)$, $x_{t4} \sim \mathcal{U}(-0.5, 0.5)$ e $z_t \sim \mathcal{U}(0.5, 1.5)$ e mantidas fixas para cada réplica de Monte Carlo. Foram considerados três diferentes cenários para a resposta média: $\mu_t \in (0.02, 0.32)$ com $\beta = (-2.4, 1.2, -1.5, -1.7)^\top$; $\mu_t \in (0.19, 0.86)$ com $\beta = (-1.7, -1.8, 1.2, -1.3)^\top$ e $\mu_t \in (0.78, 0.98)$ com $\beta = (2.1, -1.5, -1.6, -1.2)^\top$. Além disso, para medir a intensidade de dispersão não constante, consideramos $\lambda = \sigma_{\max}^2 / \sigma_{\min}^2$ e reportamos os resultados para $\lambda \approx 12$ com $\gamma = (-1.3, -1.6)^\top$; $\lambda \approx 45$ com $\gamma = (-1.3, -2.1)^\top$ e $\lambda \approx 128$ com $\gamma = (-1.3, -2.4)^\top$. Os tamanhos amostrais considerados foram $n = 40, 80$ e 120 . Para os dois últimos casos foram geradas inicialmente $n = 40$ observações das covariáveis, sendo estas replicadas duas e três vezes, respectivamente, para obter os tamanhos amostrais $n = 80$ e $n = 120$. Isso foi feito para garantir que a intensidade de dispersão não constante fosse a mesma para todos os tamanhos amostrais. O número de réplicas de Monte Carlo foi $R = 10000$ e o número de réplicas bootstrap, $B = 500$. As estimativas dos parâmetros em (3.3) foram obtidas por maximização da função de log-verossimilhança usando o método de otimização não linear Escore de Fisher.

Para cada estimativa de $\hat{\theta}$ foram geradas B réplicas bootstrap de forma paramétrica e calculada a estimativa corrigida bootstrap. Além disso, para cada réplica de Monte Carlo e para

cada estimativa de máxima verossimilhança dos parâmetros do modelo foram obtidos intervalos de confiança do tipo assintótico, bootstrap percentil e bootstrap- t . Todos os intervalos foram obtidos estimando os dois limites separadamente.

Para analisar os resultados da estimação pontual foram calculados para cada tamanho amostral e cada estimador o viés relativo e a raiz do erro quadrático médio. No que se refere à estimação intervalar calculamos as coberturas empíricas dos intervalos, obtidas a partir das frequências relativas em que os intervalos contiveram o verdadeiro valor do parâmetro, o limite inferior (superior) médio, o comprimento médio e os percentuais de vezes em que o limite inferior (superior) foi maior (menor) que o verdadeiro valor do parâmetro.

As avaliações numéricas foram realizadas utilizando a linguagem de programação Ox (DOORNIK, 2001; DOORNIK, 2006; DOORNIK, 2009). Os resultados gráficos foram obtidos utilizando o ambiente de programação R (CRIBARI-NETO; ZARKOS, 1999; DALGAARD, 2002; VENABLES; RIPLEY, 2002).

3.6.1 Resultados da Avaliação Numérica da Estimação Pontual

Nas Tabelas 1, 2 e 3 consideramos, respectivamente, os cenários em que $\mu_t \in (0.02, 0.32)$ ($\mu_t \approx 0$), $\mu_t \in (0.19, 0.86)$ ($\mu_t \approx 0.5$) e $\mu_t \in (0.78, 0.98)$ ($\mu_t \approx 1$). Nestas tabelas são apresentados os vieses relativos e as raízes quadradas dos Erros Quadráticos Médios (EQM) dos estimadores dos parâmetros para $n = 40, 80$ e 120 e $\lambda \approx 12, 45$ e 128 . Para $\theta = (\theta_1, \theta_2, \theta_3, \theta_4, \theta_5, \theta_6)^\top = (\beta_1, \beta_2, \beta_3, \beta_4, \gamma_1, \gamma_2)^\top$, tem-se que

$$\bar{\theta}_m = \frac{1}{R} \sum_{r=1}^R \hat{\theta}_m^{(r)}, \quad \widehat{\text{Viés}}(\hat{\theta}_m) = \bar{\theta}_m - \theta_m, \quad \widehat{\text{Viés}}_{rel} = \frac{\bar{\theta}_m - \theta_m}{\theta_m} \quad \text{e}$$

$$\sqrt{\text{EQM}} = \sqrt{\frac{1}{R} \sum_{r=1}^R (\hat{\theta}_m^{(r)} - \theta_m)^2}, \quad m = 1, \dots, 6,$$

em que o $\widehat{\text{Viés}}_{rel}$ e $\sqrt{\text{EQM}}$ são o viés relativo e a raiz do EQM dos estimadores do modelo. Adicionalmente, $\hat{\theta}_m^{(r)}$ representa a estimativa de θ_m obtida na r -ésima réplica de Monte Carlo, $r = 1, \dots, 10000$.

Observamos na Tabela 1 que, em módulo, as estimativas do viés relativo dos estimadores corrigidos bootstrap foram consideravelmente menores que as dos estimadores de máxima verossimilhança. Por exemplo, para $n = 40$ e $\lambda \approx 45$, a estimativa do viés relativo do estimador $\hat{\beta}_3$ (BOOT) foi 0.0003, enquanto que a do estimador $\hat{\beta}_3$ (EMV) 0.0012.

Na Tabela 2 encontram-se os resultados de simulação para o caso em que $\mu_t \in (0.19, 0.86)$, ou seja, $\mu_t \approx 0.5$. Observamos que, de forma análoga à Tabela 1, os estimadores corrigidos bootstrap apresentam melhores desempenhos em termos do viés relativo. Esse desempenho é superior ao apresentado para o cenário em que $\mu_t \in (0.02, 0.32)$, isto é, $\mu_t \approx 0$. Além disso, verificamos que as estimativas da $\sqrt{\text{EQM}}$ dos estimadores corrigidos bootstrap foram menores que as dos estimadores de máxima verossimilhança.

Na Tabela 3 apresentamos os resultados de simulação para $\mu_t \in (0.78, 0.98)$, ou seja, $\mu_t \approx 1$. Assim como acontecem nas Tabelas 1 e 2, em módulo as estimativas do viés relativo dos estimadores corrigidos bootstrap foram menores que as dos estimadores de máxima verossimilhança. Além disso, percebemos nas Tabelas 1, 2 e 3, que as estimativas da $\sqrt{\text{EQM}}$ de todos os estimadores diminuem quando o tamanho da amostra cresce.

Portanto, de modo geral, os estimadores de máxima verossimilhança e os corrigidos via esquema bootstrap apresentam bom desempenho em pequenas amostras. Os estimadores corrigidos apresentaram vieses inferiores aos dos estimadores de máxima verossimilhança para os diferentes cenários analisados, evidenciando a efetividade do esquema bootstrap na correção do viés. No entanto, essa superioridade fica mais evidente para o caso em que $\mu_t \in (0.19, 0.86)$, uma vez que nesse cenário o estimador bootstrap apresenta menores vieses para os parâmetros do modelo e para os diferentes níveis de dispersão não constante.

Os estimadores de máxima verossimilhança dos parâmetros do submodelo da dispersão, em geral, tendem a ser mais viesados do que os estimadores dos parâmetros do submodelo da média, especialmente, para o parâmetro γ_1 . Por exemplo, para $\mu_t \in (0.02, 0.32)$, $n = 120$ e $\lambda \approx 45$, a estimativa do viés relativo do estimador $\hat{\gamma}_1$ foi 0.0429, enquanto que a do estimador $\hat{\beta}_1$ 0.0006. Além disso, temos que as estimativas do viés relativo dos estimadores corrigidos $\bar{\gamma}_1$ e $\bar{\beta}_1$ foram 0.0009 e 0.0001, respectivamente. Note que para o parâmetro γ_1 em especial, o estimador corrigido bootstrap apresenta uma substancial redução do viés quando comparado a sua contrapartida não corrigida. Reforçando assim a importância da utilização do esquema proposto na correção de viés de estimadores do modelo de regressão simplex não linear.

Tabela 1 – Vieses relativos e raízes dos erros quadráticos médios dos EMVs e dos EMVs corrigidos bootstrap (Boot) dos parâmetros do Modelo (3.3) para $\mu_t \approx 0$.

$\mu_t \in (0.02, 0.32)$													
n	Parâm.	$\lambda \approx 12$				$\lambda \approx 45$				$\lambda \approx 128$			
		EMV		BOOT		EMV		BOOT		EMV		BOOT	
		Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$
40	β_1	0.0003	0.0819	-0.0008	0.0819	0.0012	0.0853	0.0002	0.0854	0.0016	0.0894	0.0005	0.0895
	β_2	-0.0009	0.1306	0.0007	0.1315	0.0040	0.1340	0.0039	0.1351	0.0036	0.1431	0.0021	0.1423
	β_3	0.0028	0.1375	0.0019	0.1376	0.0012	0.1428	0.0003	0.1431	0.0013	0.1528	-0.0002	0.1530
	β_4	-0.0007	0.1149	-0.0006	0.1152	-0.0007	0.1156	-0.0005	0.1157	0.0006	0.1228	0.0009	0.1230
	γ_1	0.1363	0.2739	-0.0015	0.2391	0.1363	0.2731	-0.0006	0.2368	0.1381	0.2706	0.0028	0.2331
	γ_2	-0.0097	0.3233	0.0011	0.3106	0.0001	0.2327	0.0008	0.2287	0.0007	0.1913	0.0008	0.1878
80	β_1	0.0012	0.0568	0.0006	0.0569	0.0001	0.0588	-0.0005	0.0590	0.0008	0.0638	0.0003	0.0639
	β_2	-0.0010	0.0907	0.0001	0.0913	0.0016	0.0936	0.0007	0.0947	0.0032	0.0988	-0.0006	0.0981
	β_3	-0.0006	0.0933	-0.0011	0.0934	0.0018	0.0982	0.0014	0.0983	0.0001	0.1076	-0.0006	0.1077
	β_4	-0.0004	0.0791	-0.0004	0.0792	-0.0002	0.0806	0.0001	0.0807	-0.0014	0.0853	-0.0009	0.0855
	γ_1	0.0673	0.1734	0.0019	0.1602	0.0653	0.1735	0.0012	0.1621	0.0603	0.1684	-0.0010	0.1582
	γ_2	-0.0004	0.2087	0.0030	0.2059	-0.0004	0.1538	0.0008	0.1543	-0.0029	0.1283	-0.0003	0.1274
120	β_1	0.0007	0.0465	0.0003	0.0465	0.0006	0.0495	0.0001	0.0495	0.0002	0.0524	-0.0002	0.0524
	β_2	-0.0005	0.0749	0.0003	0.0753	0.0007	0.0770	-0.0002	0.0780	0.0061	0.0815	0.0018	0.0811
	β_3	-0.0003	0.0769	-0.0006	0.0769	-0.0003	0.0813	-0.0005	0.0814	0.0014	0.0886	0.0010	0.0887
	β_4	-0.0002	0.0651	-0.0002	0.0652	-0.0008	0.0659	-0.0007	0.0659	-0.0003	0.0703	0.0002	0.0706
	γ_1	0.0430	0.1344	0.0001	0.1280	0.0429	0.1348	0.0009	0.1284	0.0390	0.1327	-0.0001	0.1274
	γ_2	-0.0014	0.1653	0.0004	0.1638	-0.0008	0.1232	0.0001	0.1238	-0.0035	0.1046	-0.0008	0.1041

Fonte: próprio autor.

Tabela 2 – Vieses relativos e raízes dos erros quadráticos médios dos EMVs e dos EMVs corrigidos bootstrap (Boot) dos parâmetros do Modelo (3.3) para $\mu_t \approx 0.5$.

$\mu_t \in (0.19, 0.86)$													
n	Parâm.	$\lambda \approx 12$				$\lambda \approx 45$				$\lambda \approx 128$			
		EMV		BOOT		EMV		BOOT		EMV		BOOT	
		Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$
40	β_1	0.0022	0.1189	0.0005	0.1187	0.0017	0.1246	0.0001	0.1243	0.0029	0.1260	0.0011	0.1256
	β_2	0.0012	0.0977	-0.0004	0.0975	0.0015	0.1016	-0.0001	0.1013	0.0011	0.1018	-0.0007	0.1015
	β_3	0.0072	0.2201	0.0040	0.2194	0.0040	0.2324	0.0007	0.2317	0.0067	0.2367	0.0033	0.2357
	β_4	0.0049	0.1772	0.0009	0.1768	0.0075	0.1759	0.0030	0.1754	0.0042	0.1805	-0.0004	0.1800
	γ_1	0.1313	0.2739	-0.0037	0.2410	0.1392	0.2746	0.0014	0.2365	0.1374	0.2731	-0.0015	0.2349
	γ_2	-0.0199	0.3328	0.0003	0.3095	0.0024	0.2346	0.0026	0.2248	0.0041	0.1926	0.0007	0.1880
80	β_1	0.0011	0.0834	0.0002	0.0833	0.0008	0.0866	-0.0001	0.0865	0.0007	0.0882	-0.0003	0.0881
	β_2	0.0007	0.0685	-0.0004	0.0684	0.0003	0.0712	-0.0007	0.0712	0.0011	0.0716	0.0001	0.0715
	β_3	0.0012	0.1545	-0.0005	0.1543	0.0021	0.1623	0.0003	0.1619	0.0021	0.1661	0.0003	0.1657
	β_4	0.0024	0.1236	0.0002	0.1234	0.0028	0.1268	0.0004	0.1264	0.0028	0.1269	0.0002	0.1268
	γ_1	0.0636	0.1711	-0.0001	0.1586	0.0652	0.1710	0.0001	0.1585	0.0656	0.1704	0.0001	0.1570
	γ_2	-0.0047	0.2084	0.0028	0.2021	0.0004	0.1540	0.0005	0.1517	0.0017	0.1271	0.0006	0.1260
120	β_1	0.0007	0.0677	0.0001	0.0676	0.0006	0.0703	-0.0001	0.0702	0.0003	0.0719	-0.0005	0.0719
	β_2	0.0007	0.0555	0.0001	0.0555	0.0006	0.0580	-0.0001	0.0580	0.0006	0.0579	-0.0002	0.0578
	β_3	0.0005	0.1262	-0.0007	0.1262	0.0002	0.1323	-0.0010	0.1322	0.0014	0.1342	0.0001	0.1340
	β_4	0.0021	0.1001	0.0006	0.1001	0.0025	0.1021	0.0008	0.1021	0.0015	0.1036	-0.0003	0.1036
	γ_1	0.0416	0.1338	-0.0001	0.1271	0.0427	0.1343	0.0001	0.1273	0.0432	0.1342	0.0003	0.1271
	γ_2	-0.0042	0.1683	0.0004	0.1653	-0.0002	0.1245	-0.0001	0.1235	0.0009	0.1040	0.0003	0.1036

Fonte: próprio autor.

Tabela 3 – Vieses relativos e raízes dos erros quadráticos médios dos EMVs e dos EMVs corrigidos bootstrap (Boot) dos parâmetros do Modelo (3.3) para $\mu_t \approx 1$.

$\mu_t \in (0.78, 0.98)$													
n	Parâm.	$\lambda \approx 12$				$\lambda \approx 45$				$\lambda \approx 128$			
		EMV		BOOT		EMV		BOOT		EMV		BOOT	
		Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$	Viés Rel.	$\sqrt{EQM}$
40	β_1	0.0009	0.0632	0.0003	0.0632	0.0007	0.0666	-0.0001	0.0668	0.0007	0.0671	-0.0001	0.0672
	β_2	-0.0003	0.0560	0.0001	0.0560	-0.0003	0.0592	0.0001	0.0593	-0.0006	0.0617	-0.0002	0.0617
	β_3	0.0004	0.1228	0.0007	0.1229	0.0001	0.1292	0.0003	0.1293	-0.0011	0.1317	-0.0010	0.1317
	β_4	-0.0016	0.0968	-0.0013	0.0970	0.0004	0.1013	0.0009	0.1015	0.0003	0.1010	0.0009	0.1011
	γ_1	0.1301	0.2709	-0.0036	0.2383	0.1373	0.2730	-0.0005	0.2358	0.1384	0.2737	-0.0014	0.2367
	γ_2	-0.0229	0.3358	0.0021	0.3096	-0.0004	0.2394	-0.0001	0.2276	0.0042	0.1911	0.0004	0.1861
80	β_1	0.0004	0.0442	0.0001	0.0442	0.0004	0.0473	0.0001	0.0473	0.0008	0.0478	0.0004	0.0478
	β_2	0.0001	0.0390	0.0001	0.0390	-0.0006	0.0414	-0.0004	0.0414	-0.0010	0.0433	-0.0008	0.0434
	β_3	0.0005	0.0860	0.0006	0.0861	-0.0011	0.0909	-0.0011	0.0910	-0.0003	0.0931	-0.0002	0.0932
	β_4	0.0007	0.0694	0.0009	0.0695	0.0004	0.0707	0.0007	0.0708	0.0001	0.0716	0.0005	0.0717
	γ_1	0.0645	0.1711	0.0014	0.1593	0.0670	0.1721	0.0015	0.1576	0.0667	0.1720	0.0004	0.1586
	γ_2	-0.0085	0.2148	0.0005	0.2074	0.0014	0.1556	0.0011	0.1528	0.0029	0.1291	0.0011	0.1278
120	β_1	0.0002	0.0368	-0.0001	0.0368	0.0005	0.0383	0.0002	0.0383	0.0003	0.0387	0.0001	0.0387
	β_2	0.0003	0.0314	0.0004	0.0314	0.0001	0.0337	0.0001	0.0338	-0.0001	0.0352	0.0001	0.0352
	β_3	-0.0001	0.0712	0.0001	0.0713	0.0003	0.0754	0.0004	0.0754	0.0001	0.0752	0.0001	0.0752
	β_4	-0.0013	0.0559	-0.0012	0.0560	-0.0006	0.0574	-0.0004	0.0574	-0.0003	0.0586	-0.0001	0.0587
	γ_1	0.0422	0.1340	0.0008	0.1275	0.0456	0.1359	0.0027	0.1277	0.0446	0.1344	0.0010	0.1276
	γ_2	-0.0040	0.1677	0.0014	0.1644	0.0009	0.1228	0.0007	0.1215	0.0020	0.1022	0.0009	0.1016

Fonte: próprio autor.

3.6.2 Resultados Numéricos sobre Intervalos de Confiança

Avaliamos também os comportamentos dos estimadores intervalares descritos nas seções anteriores. Apresentamos resultados para os níveis nominais de cobertura $1 - \alpha$ ($0 < \alpha < 1$) iguais a 0.90, 0.95 e 0.99. Todas as estimativas intervalares foram obtidas com base em estimadores intervalares com probabilidades $1 - \alpha$ de incluir o verdadeiro valor de θ e $\alpha/2$ do limite inferior (superior) ser maior (menor) que θ . As Tabelas 4, 5 e 6 consideram probabilidades nominais de cobertura iguais a 0.90, 0.95 e 0.99, respectivamente. Nestas tabelas são apresentadas as taxas de cobertura dos intervalos de confiança assintótico (ICA), bootstrap- t (ICBT) e bootstrap Percentil (ICBP) para os estimadores do Modelo (3.3).

Em relação à Tabela 4, notamos que o intervalo de confiança assintótico apresenta baixas probabilidades empíricas de cobertura em pequenas amostras. Esse comportamento é similar para todos os cenários da resposta média e para as intensidades de dispersão não constante $\lambda \approx 12$, $\lambda \approx 45$ e $\lambda \approx 128$.

De modo geral, o intervalo de confiança bootstrap percentil apresenta desempenho superior apenas para o parâmetro γ_2 . Para o parâmetro de dispersão γ_1 , o intervalo ICPE apresenta baixa taxa de cobertura. Por exemplo, para $\mu_t \in (0.02, 0.32)$, $\lambda \approx 45$ e $n = 40$, a taxa foi igual a 0.6633. Para os demais parâmetros, no geral o intervalo ICPE se saiu melhor que o ICA. Percebemos que para os parâmetros do submodelo da média e o cenário $\mu_t \in (0.02, 0.32)$, $\lambda \approx 12$, $\lambda \approx 128$ e $n = 120$, os intervalos de confiança bootstrap percentil e assintótico apresentam probabilidades empíricas de cobertura semelhantes e próximas de 0.89. Quando $\mu_t \in (0.19, 0.86)$ as probabilidades ficam em torno de 0.90.

O intervalo de confiança bootstrap- t é o que apresenta melhor desempenho com probabilidades empíricas de cobertura mais próximas do nível nominal 90%. Observamos ainda que a medida que o tamanho da amostra cresce as taxas de cobertura dos três intervalos se aproximam do nível nominal.

Na Tabela 5 encontram-se os resultados de simulação para o nível nominal 95%. Os intervalos ICA e ICPE apresentam baixas probabilidades empíricas de cobertura. Este resultado é mais notório para o parâmetro de dispersão γ_1 . Por exemplo, para $n = 40$ as taxas ficam próximas de 87% e 74%, respectivamente. Percebemos ainda que as taxas para os parâmetros do submodelo da média são semelhantes. Por exemplo, para $n = 40$, $n = 80$ e $n = 120$ as taxas ficam em torno de 0.91, 0.93 e 0.94, respectivamente. De modo geral, o intervalo de confiança bootstrap percentil apresenta desempenho superior apenas para o parâmetro γ_2 .

Assim como na Tabela 4, o intervalo de confiança bootstrap- t é o que apresenta melhor resultado com cobertura próxima do nível nominal 95%, além das taxas de cobertura dos três intervalos se aproximarem do nível nominal quando aumenta-se o tamanho da amostra.

Por último, a Tabela 6 apresenta os resultados de simulação para o nível nominal 99%. De forma análoga às Tabelas 4 e 5, os intervalos de confiança assintótico e bootstrap percentil apresentam desempenhos inferiores ao do intervalo de confiança bootstrap- t . O intervalo ICBT apresenta excelentes taxas de cobertura para todos os parâmetros do modelo.

Assim como ocorreu nas Tabelas 4 e 5, a pior cobertura que os intervalos de confiança ICA e ICPE apresentaram foi para o parâmetro de dispersão γ_1 . Por exemplo, para $n = 40$ as taxas ficaram próximas de 95% e 87%, respectivamente. O melhor resultado do intervalo ICPE foi para o parâmetro γ_2 . Com relação aos parâmetros do submodelo da média, os intervalos de confiança ICPE e ICA apresentaram cobertura em torno de 0.98 para $n = 120$ e em torno de 0.97 para $n = 40$. Notamos, de forma geral, que conforme aumentamos o tamanho da amostra as probabilidades empíricas de cobertura de todos os intervalos se aproximam do nível nominal.

As Tabelas 7 até 11 apresentam as médias dos limites inferiores (Inferior) e superiores (Superior), médias dos comprimentos (Tamanho) e a probabilidade empírica de cobertura (Cobertura) dos intervalos de confiança. Adicionalmente, são apresentados os percentuais de vezes que o verdadeiro valor do parâmetro foi menor que o limite inferior do IC (% Esquerdo) e foi maior que o limite superior do IC (% Direito). Essas duas quantidades avaliam o balanceamento do intervalo. O balanceamento perfeito ocorre quando esses dois percentuais são idênticos. Essas propriedades dos estimadores intervalares foram avaliadas apenas para os parâmetros de não linearidade da média e da dispersão, β_2 e γ_2 , respectivamente.

Nas Tabelas 7, 8 e 9 são apresentados resultados para β_2 , $n = 40$, $n = 80$ e $n = 120$, e apenas para $\mu_t \in (0.02, 0.32)$, em que $\beta_2 = 1.2$. Note que se $\beta_2 = 1$ o modelo passa a ser linear. Para $n = 40$, $\lambda \approx 12$ e $\lambda \approx 45$ apenas o intervalo ICBT e para $1 - \alpha = 0.99$, considera essa possibilidade. Isso ocorre porque o intervalo de confiança bootstrap- t é o que apresenta o maior comprimento médio entre os três intervalos. Sua probabilidade empírica de cobertura é a mais próxima da verdadeira, no entanto, o comprimento médio do intervalo é maior, o que permite a inclusão de $\beta_2 = 1.0$, ou seja, linearidade (Tabela 7). No entanto, esse comportamento do intervalo ICBT não se mantém para $n = 80$ e $n = 120$, Tabelas 8 e 9, respectivamente.

Quando $\lambda \approx 128$ todos intervalos, para todos os níveis de confiança, incluem a possibilidade

Tabela 4 – Taxas de cobertura dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para os parâmetros do modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $1 - \alpha = 0.90$.

$\mu_t \in (0.02, 0.32)$										
n	Parâmetros	$\lambda \approx 12$			$\lambda \approx 45$			$\lambda \approx 128$		
		ICA	ICBT	ICPE	ICA	ICBT	ICPE	ICA	ICBT	ICPE
40	β_1	0.8524	0.8976	0.8583	0.8534	0.8969	0.8578	0.8595	0.8975	0.8617
	β_2	0.8593	0.9016	0.8677	0.8582	0.8962	0.8621	0.8558	0.8987	0.8550
	β_3	0.8481	0.8927	0.8552	0.8492	0.8942	0.8547	0.8617	0.8996	0.8638
	β_4	0.8535	0.8921	0.8569	0.8532	0.8919	0.8558	0.8575	0.8961	0.8591
	γ_1	0.8003	0.8944	0.6633	0.8043	0.9018	0.6633	0.8077	0.9051	0.6585
	γ_2	0.8533	0.8939	0.8893	0.8610	0.8908	0.9020	0.8632	0.8925	0.8970
80	β_1	0.8813	0.8997	0.8813	0.8805	0.8986	0.8847	0.8819	0.8991	0.8810
	β_2	0.8858	0.9011	0.8865	0.8823	0.8952	0.8846	0.8838	0.9034	0.8815
	β_3	0.8832	0.9023	0.8853	0.8826	0.9011	0.8844	0.8835	0.9002	0.8844
	β_4	0.8797	0.8954	0.8804	0.8815	0.8976	0.8824	0.8823	0.8972	0.8808
	γ_1	0.8522	0.8961	0.7684	0.8473	0.8920	0.7697	0.8606	0.8966	0.7908
	γ_2	0.8792	0.8893	0.8989	0.8808	0.8866	0.9020	0.8840	0.8930	0.8966
120	β_1	0.8828	0.8980	0.8829	0.8756	0.8863	0.8761	0.8874	0.8981	0.8869
	β_2	0.8899	0.8997	0.8892	0.8899	0.8956	0.8921	0.8877	0.8998	0.8854
	β_3	0.8862	0.8960	0.8844	0.8775	0.8890	0.8768	0.8862	0.8980	0.8865
	β_4	0.8854	0.8972	0.8849	0.8896	0.9002	0.8880	0.8886	0.8991	0.8897
	γ_1	0.8718	0.8977	0.8151	0.8681	0.8963	0.8147	0.8751	0.9002	0.8267
	γ_2	0.8852	0.8929	0.8988	0.8913	0.8930	0.9042	0.8843	0.8901	0.8917
$\mu_t \in (0.19, 0.86)$										
n	Parâmetros	$\lambda \approx 12$			$\lambda \approx 45$			$\lambda \approx 128$		
		ICA	ICBT	ICPE	ICA	ICBT	ICPE	ICA	ICBT	ICPE
40	β_1	0.8538	0.8993	0.8624	0.8535	0.8983	0.8586	0.8599	0.9024	0.8619
	β_2	0.8506	0.9019	0.8560	0.8534	0.9003	0.8555	0.8615	0.9063	0.8655
	β_3	0.8553	0.8979	0.8621	0.8527	0.8984	0.8588	0.8585	0.9034	0.8611
	β_4	0.8593	0.9005	0.8617	0.8605	0.9006	0.8633	0.8614	0.8995	0.8636
	γ_1	0.8032	0.8956	0.6727	0.8007	0.9016	0.6598	0.8030	0.9040	0.6577
	γ_2	0.8478	0.9056	0.8860	0.8529	0.8954	0.8908	0.8609	0.8967	0.8983
80	β_1	0.8769	0.8962	0.8793	0.8797	0.9010	0.8803	0.8786	0.8992	0.8805
	β_2	0.8731	0.8981	0.8759	0.8762	0.8996	0.8755	0.8830	0.9037	0.8839
	β_3	0.8748	0.8957	0.8772	0.8789	0.8977	0.8777	0.8810	0.9008	0.8829
	β_4	0.8803	0.8981	0.8789	0.8788	0.8980	0.8777	0.8843	0.9020	0.8814
	γ_1	0.8576	0.9034	0.7813	0.8555	0.9023	0.7755	0.8561	0.9038	0.7758
	γ_2	0.8792	0.8982	0.8994	0.8818	0.8998	0.9002	0.8838	0.9008	0.9026
120	β_1	0.8846	0.9002	0.8848	0.8860	0.8996	0.8894	0.8852	0.8978	0.8864
	β_2	0.8802	0.8992	0.8821	0.8838	0.8972	0.8826	0.8891	0.9038	0.8898
	β_3	0.8850	0.8982	0.8872	0.8859	0.8967	0.8855	0.8864	0.8981	0.8860
	β_4	0.8882	0.8986	0.8880	0.8906	0.9030	0.8898	0.8886	0.9017	0.8897
	γ_1	0.8727	0.9041	0.8203	0.8704	0.9004	0.8161	0.8735	0.9026	0.8108
	γ_2	0.8848	0.8948	0.8940	0.8855	0.8948	0.8968	0.8871	0.8968	0.8983
$\mu_t \in (0.78, 0.98)$										
n	Parâmetros	$\lambda \approx 12$			$\lambda \approx 45$			$\lambda \approx 128$		
		ICA	ICBT	ICPE	ICA	ICBT	ICPE	ICA	ICBT	ICPE
40	β_1	0.8570	0.9000	0.8634	0.8530	0.8941	0.8551	0.8585	0.8981	0.8597
	β_2	0.8481	0.8967	0.8579	0.8513	0.8972	0.8589	0.8587	0.9008	0.8624
	β_3	0.8586	0.8983	0.8643	0.8593	0.9019	0.8635	0.8582	0.8962	0.8612
	β_4	0.8597	0.9005	0.8660	0.8556	0.8973	0.8583	0.8627	0.9001	0.8662
	γ_1	0.8136	0.8985	0.6812	0.8043	0.9012	0.6642	0.8036	0.8972	0.6567
	γ_2	0.8463	0.9063	0.8861	0.8514	0.8939	0.8872	0.8589	0.8987	0.8936
80	β_1	0.8821	0.9017	0.8825	0.8757	0.8967	0.8754	0.8805	0.8985	0.8782
	β_2	0.8779	0.8971	0.8806	0.8789	0.8985	0.8817	0.8799	0.8991	0.8796
	β_3	0.8774	0.8941	0.8791	0.8770	0.8943	0.8777	0.8804	0.8966	0.8799
	β_4	0.8771	0.8951	0.8775	0.8810	0.8973	0.8800	0.8747	0.8902	0.8754
	γ_1	0.8501	0.8995	0.7789	0.8499	0.9029	0.7698	0.8512	0.8986	0.7690
	γ_2	0.8714	0.8999	0.8898	0.8795	0.8987	0.8977	0.8809	0.8970	0.8977
120	β_1	0.8816	0.8947	0.8832	0.8840	0.8947	0.8850	0.8913	0.9009	0.8873
	β_2	0.8898	0.9011	0.8923	0.8928	0.9031	0.8923	0.8846	0.8948	0.8847
	β_3	0.8856	0.8956	0.8855	0.8789	0.8898	0.8795	0.8873	0.8994	0.8862
	β_4	0.8842	0.8958	0.8839	0.8868	0.8987	0.8863	0.8841	0.8977	0.8843
	γ_1	0.8707	0.9042	0.8164	0.8709	0.9005	0.8080	0.8653	0.9001	0.8070
	γ_2	0.8876	0.9019	0.9019	0.8872	0.9011	0.8981	0.8830	0.8951	0.8940

Fonte: próprio autor.

Tabela 5 – Taxas de cobertura dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para os parâmetros do modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $1 - \alpha = 0.95$.

$\mu_t \in (0.02, 0.32)$										
n	Parâmetros	$\lambda \approx 12$			$\lambda \approx 45$			$\lambda \approx 128$		
		ICA	ICBT	ICPE	ICA	ICBT	ICPE	ICA	ICBT	ICPE
40	β_1	0.9131	0.9479	0.9187	0.9147	0.9490	0.9177	0.9185	0.9493	0.9200
	β_2	0.9190	0.9499	0.9233	0.9187	0.9498	0.9216	0.9166	0.9495	0.9159
	β_3	0.9088	0.9443	0.9143	0.9139	0.9473	0.9158	0.9173	0.9478	0.9183
	β_4	0.9111	0.9378	0.9118	0.9139	0.9422	0.9140	0.9176	0.9451	0.9163
	γ_1	0.8720	0.9455	0.7437	0.8720	0.9513	0.7437	0.8760	0.9514	0.7448
	γ_2	0.9147	0.9540	0.9441	0.9215	0.9392	0.9498	0.9227	0.9396	0.9482
80	β_1	0.9337	0.9490	0.9343	0.9335	0.9481	0.9340	0.9370	0.9510	0.9368
	β_2	0.9375	0.9509	0.9376	0.9346	0.9457	0.9374	0.9349	0.9523	0.9336
	β_3	0.9365	0.9515	0.9367	0.9363	0.9511	0.9372	0.9374	0.9476	0.9373
	β_4	0.9328	0.9475	0.9338	0.9347	0.9445	0.9332	0.9372	0.9484	0.9359
	γ_1	0.9145	0.9483	0.8399	0.9074	0.9440	0.8377	0.9173	0.9468	0.8602
	γ_2	0.9360	0.9421	0.9498	0.9369	0.9400	0.9485	0.9407	0.9422	0.9464
120	β_1	0.9386	0.9462	0.9372	0.9310	0.9411	0.9316	0.9424	0.9498	0.9407
	β_2	0.9414	0.9478	0.9404	0.9422	0.9463	0.9418	0.9402	0.9513	0.9392
	β_3	0.9380	0.9468	0.9364	0.9337	0.9443	0.9341	0.9424	0.9497	0.9412
	β_4	0.9408	0.9493	0.9401	0.9421	0.9470	0.9389	0.9406	0.9467	0.9412
	γ_1	0.9277	0.9458	0.8791	0.9233	0.9479	0.8779	0.9312	0.9504	0.8890
	γ_2	0.9395	0.9428	0.9484	0.9425	0.9436	0.9503	0.9403	0.9425	0.9457
$\mu_t \in (0.19, 0.86)$										
n	Parâmetros	$\lambda \approx 12$			$\lambda \approx 45$			$\lambda \approx 128$		
		ICA	ICBT	ICPE	ICA	ICBT	ICPE	ICA	ICBT	ICPE
40	β_1	0.9133	0.9482	0.9181	0.9139	0.9473	0.9168	0.9162	0.9505	0.9172
	β_2	0.9092	0.9503	0.9155	0.9138	0.9515	0.9185	0.9207	0.9515	0.9234
	β_3	0.9137	0.9508	0.9184	0.9126	0.9477	0.9145	0.9182	0.9513	0.9200
	β_4	0.9155	0.9457	0.9189	0.9192	0.9532	0.9213	0.9174	0.9500	0.9201
	γ_1	0.8720	0.9460	0.7541	0.8667	0.9515	0.7430	0.8706	0.9487	0.7403
	γ_2	0.9151	0.9640	0.9432	0.9148	0.9427	0.9459	0.9205	0.9434	0.9484
80	β_1	0.9342	0.9497	0.9336	0.9342	0.9498	0.9351	0.9333	0.9502	0.9350
	β_2	0.9297	0.9483	0.9325	0.9318	0.9482	0.9324	0.9365	0.9525	0.9368
	β_3	0.9305	0.9477	0.9313	0.9305	0.9449	0.9329	0.9370	0.9500	0.9356
	β_4	0.9344	0.9480	0.9336	0.9333	0.9485	0.9328	0.9371	0.9501	0.9352
	γ_1	0.9167	0.9484	0.8473	0.9138	0.9482	0.8449	0.9164	0.9503	0.8462
	γ_2	0.9371	0.9467	0.9490	0.9392	0.9459	0.9494	0.9423	0.9491	0.9528
120	β_1	0.9405	0.9494	0.9419	0.9398	0.9466	0.9389	0.9408	0.9495	0.9419
	β_2	0.9374	0.9471	0.9381	0.9373	0.9474	0.9369	0.9391	0.9499	0.9395
	β_3	0.9398	0.9480	0.9378	0.9388	0.9501	0.9369	0.9386	0.9475	0.9386
	β_4	0.9414	0.9501	0.9413	0.9412	0.9500	0.9411	0.9428	0.9516	0.9420
	γ_1	0.9273	0.9507	0.8824	0.9267	0.9502	0.8806	0.9296	0.9477	0.8803
	γ_2	0.9375	0.9421	0.9463	0.9390	0.9451	0.9471	0.9400	0.9457	0.9487
$\mu_t \in (0.78, 0.98)$										
n	Parâmetros	$\lambda \approx 12$			$\lambda \approx 45$			$\lambda \approx 128$		
		ICA	ICBT	ICPE	ICA	ICBT	ICPE	ICA	ICBT	ICPE
40	β_1	0.9153	0.9515	0.9211	0.9126	0.9463	0.9163	0.9170	0.9478	0.9177
	β_2	0.9073	0.9450	0.9160	0.9125	0.9468	0.9187	0.9173	0.9475	0.9178
	β_3	0.9185	0.9502	0.9215	0.9211	0.9503	0.9232	0.9143	0.9442	0.9173
	β_4	0.9174	0.9482	0.9205	0.9178	0.9495	0.9204	0.9212	0.9498	0.9205
	γ_1	0.8731	0.9456	0.7594	0.8731	0.9495	0.7491	0.8678	0.9487	0.7430
	γ_2	0.9130	0.9613	0.9422	0.9154	0.9407	0.9453	0.9173	0.9454	0.9476
80	β_1	0.9375	0.9501	0.9371	0.9327	0.9462	0.9319	0.9322	0.9462	0.9326
	β_2	0.9329	0.9498	0.9360	0.9345	0.9489	0.9357	0.9355	0.9472	0.9341
	β_3	0.9354	0.9486	0.9364	0.9341	0.9478	0.9347	0.9376	0.9482	0.9358
	β_4	0.9312	0.9435	0.9318	0.9322	0.9455	0.9326	0.9304	0.9433	0.9289
	γ_1	0.9125	0.9474	0.8475	0.9136	0.9500	0.8415	0.9152	0.9487	0.8414
	γ_2	0.9324	0.9428	0.9433	0.9357	0.9460	0.9494	0.9336	0.9443	0.9458
120	β_1	0.9370	0.9472	0.9377	0.9370	0.9471	0.9362	0.9412	0.9492	0.9426
	β_2	0.9420	0.9506	0.9419	0.9411	0.9500	0.9419	0.9400	0.9481	0.9398
	β_3	0.9358	0.9435	0.9373	0.9331	0.9419	0.9336	0.9384	0.9474	0.9385
	β_4	0.9380	0.9453	0.9388	0.9410	0.9480	0.9389	0.9413	0.9493	0.9394
	γ_1	0.9275	0.9509	0.8808	0.9272	0.9504	0.8754	0.9252	0.9473	0.8720
	γ_2	0.9404	0.9476	0.9493	0.9405	0.9471	0.9493	0.9381	0.9449	0.9473

Fonte: próprio autor.

Tabela 6 – Taxas de cobertura dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para os parâmetros do modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $1 - \alpha = 0.99$.

$\mu_t \in (0.02, 0.32)$										
n	Parâmetros	$\lambda \approx 12$			$\lambda \approx 45$			$\lambda \approx 128$		
		ICA	ICBT	ICPE	ICA	ICBT	ICPE	ICA	ICBT	ICPE
40	β_1	0.9736	0.9881	0.9744	0.9733	0.9880	0.9731	0.9765	0.9894	0.9746
	β_2	0.9757	0.9909	0.9761	0.9767	0.9880	0.9766	0.9732	0.9885	0.9733
	β_3	0.9714	0.9879	0.9734	0.9739	0.9885	0.9728	0.9743	0.9886	0.9739
	β_4	0.9698	0.9857	0.9707	0.9712	0.9861	0.9698	0.9754	0.9877	0.9730
	γ_1	0.9480	0.9874	0.8664	0.9529	0.9891	0.8661	0.9523	0.9894	0.8687
	γ_2	0.9693	0.9887	0.9873	0.9773	0.9892	0.9902	0.9791	0.9804	0.9875
80	β_1	0.9831	0.9870	0.9810	0.9830	0.9882	0.9816	0.9847	0.9891	0.9831
	β_2	0.9853	0.9885	0.9838	0.9829	0.9879	0.9812	0.9827	0.9883	0.9808
	β_3	0.9853	0.9883	0.9845	0.9847	0.9895	0.9825	0.9842	0.9879	0.9823
	β_4	0.9837	0.9884	0.9820	0.9825	0.9873	0.9811	0.9842	0.9882	0.9815
	γ_1	0.9729	0.9890	0.9289	0.9690	0.9881	0.9268	0.9731	0.9870	0.9379
	γ_2	0.9822	0.9864	0.9894	0.9862	0.9843	0.9881	0.9864	0.9845	0.9848
120	β_1	0.9839	0.9880	0.9835	0.9839	0.9871	0.9820	0.9856	0.9860	0.9833
	β_2	0.9863	0.9887	0.9859	0.9865	0.9882	0.9848	0.9864	0.9884	0.9845
	β_3	0.9849	0.9876	0.9840	0.9852	0.9875	0.9851	0.9871	0.9886	0.9840
	β_4	0.9850	0.9871	0.9827	0.9854	0.9862	0.9844	0.9861	0.9880	0.9838
	γ_1	0.9777	0.9872	0.9504	0.9795	0.9888	0.9500	0.9817	0.9880	0.9590
	γ_2	0.9866	0.9847	0.9883	0.9871	0.9838	0.9865	0.9873	0.9850	0.9848
$\mu_t \in (0.19, 0.86)$										
n	Parâmetros	$\lambda \approx 12$			$\lambda \approx 45$			$\lambda \approx 128$		
		ICA	ICBT	ICPE	ICA	ICBT	ICPE	ICA	ICBT	ICPE
40	β_1	0.9723	0.9884	0.9735	0.9714	0.9866	0.9714	0.9753	0.9889	0.9743
	β_2	0.9676	0.9896	0.9713	0.9698	0.9896	0.9721	0.9710	0.9889	0.9732
	β_3	0.9733	0.9871	0.9735	0.9721	0.9881	0.9723	0.9746	0.9891	0.9728
	β_4	0.9699	0.9872	0.9699	0.9765	0.9898	0.9754	0.9745	0.9876	0.9723
	γ_1	0.9482	0.9875	0.8701	0.9478	0.9899	0.8631	0.9513	0.9883	0.8676
	γ_2	0.9715	0.9915	0.9866	0.9754	0.9915	0.9880	0.9772	0.9830	0.9879
80	β_1	0.9826	0.9874	0.9807	0.9820	0.9884	0.9824	0.9838	0.9884	0.9821
	β_2	0.9806	0.9874	0.9804	0.9816	0.9875	0.9805	0.9829	0.9898	0.9817
	β_3	0.9817	0.9881	0.9805	0.9819	0.9872	0.9799	0.9854	0.9902	0.9828
	β_4	0.9802	0.9856	0.9798	0.9810	0.9869	0.9792	0.9848	0.9903	0.9835
	γ_1	0.9733	0.9880	0.9366	0.9724	0.9869	0.9316	0.9743	0.9870	0.9341
	γ_2	0.9844	0.9885	0.9877	0.9848	0.9841	0.9867	0.9855	0.9879	0.9893
120	β_1	0.9863	0.9881	0.9852	0.9854	0.9885	0.9849	0.9850	0.9882	0.9842
	β_2	0.9843	0.9883	0.9825	0.9840	0.9890	0.9821	0.9840	0.9885	0.9835
	β_3	0.9856	0.9885	0.9845	0.9859	0.9892	0.9839	0.9856	0.9876	0.9830
	β_4	0.9861	0.9877	0.9845	0.9853	0.9880	0.9842	0.9876	0.9894	0.9855
	γ_1	0.9814	0.9894	0.9544	0.9807	0.9901	0.9518	0.9790	0.9888	0.9547
	γ_2	0.9860	0.9867	0.9873	0.9860	0.9850	0.9882	0.9877	0.9856	0.9864
$\mu_t \in (0.78, 0.98)$										
n	Parâmetros	$\lambda \approx 12$			$\lambda \approx 45$			$\lambda \approx 128$		
		ICA	ICBT	ICPE	ICA	ICBT	ICPE	ICA	ICBT	ICPE
40	β_1	0.9732	0.9891	0.9742	0.9734	0.9886	0.9731	0.9755	0.9884	0.9717
	β_2	0.9690	0.9881	0.9734	0.9724	0.9872	0.9728	0.9717	0.9875	0.9726
	β_3	0.9754	0.9886	0.9752	0.9765	0.9890	0.9765	0.9739	0.9867	0.9734
	β_4	0.9740	0.9876	0.9742	0.9763	0.9866	0.9749	0.9748	0.9874	0.9745
	γ_1	0.9505	0.9871	0.8737	0.9478	0.9881	0.8677	0.9483	0.9882	0.8609
	γ_2	0.9696	0.9901	0.9872	0.9747	0.9889	0.9864	0.9765	0.9824	0.9868
80	β_1	0.9839	0.9886	0.9838	0.9842	0.9887	0.9813	0.9824	0.9880	0.9814
	β_2	0.9841	0.9881	0.9839	0.9839	0.9887	0.9834	0.9826	0.9874	0.9813
	β_3	0.9844	0.9882	0.9821	0.9844	0.9882	0.9821	0.9836	0.9880	0.9814
	β_4	0.9821	0.9855	0.9802	0.9833	0.9884	0.9819	0.9821	0.9873	0.9800
	γ_1	0.9723	0.9874	0.9321	0.9743	0.9873	0.9353	0.9718	0.9880	0.9296
	γ_2	0.9830	0.9882	0.9870	0.9840	0.9861	0.9895	0.9825	0.9846	0.9871
120	β_1	0.9848	0.9877	0.9825	0.9831	0.9855	0.9819	0.9863	0.9882	0.9847
	β_2	0.9856	0.9865	0.9843	0.9862	0.9880	0.9838	0.9869	0.9899	0.9848
	β_3	0.9853	0.9869	0.9830	0.9835	0.9846	0.9816	0.9847	0.9866	0.9815
	β_4	0.9839	0.9867	0.9833	0.9863	0.9871	0.9841	0.9857	0.9882	0.9847
	γ_1	0.9802	0.9893	0.9550	0.9790	0.9893	0.9522	0.9773	0.9865	0.9500
	γ_2	0.9853	0.9854	0.9864	0.9865	0.9864	0.9869	0.9863	0.9866	0.9880

Fonte: próprio autor.

de linearidade quando $n = 40$, ou seja, $\beta_2 = 1.0$ (Tabela 7). Esse resultado é interessante, pois mostra como a intensidade de dispersão não constante afeta negativamente os desempenhos dos três intervalos de confiança considerados. Quando o tamanho da amostra aumenta o problema é amenizado, no entanto, o desempenho do intervalo ICBT ainda é o mais afetado negativamente. Por exemplo, para $n = 80$ e $1 - \alpha = 0.90$ apenas o intervalo bootstrap- t considera a possibilidade de linearidade, ICA= (1.0065, 1.4012), ICBT= (0.9929, 1.4098) e ICPE= (1.0099, 1.4062). O desempenho do intervalo ICBT é afetado negativamente com respeito ao tamanho do intervalo que é tipicamente maior que o tamanho dos outros dois intervalos.

As Tabelas 10 e 11 apresentam os resultados de simulação para β_2 ($\mu_t \approx 0.5$ e $\mu_t \approx 1$) e γ_2 ($\forall \mu_t$), respectivamente, quando $n = 40$ e $\lambda \approx 128$. Para os cenários $\mu_t \in (0.19, 0.86)$ e $\mu_t \in (0.78, 0.98)$, β_2 é igual a -1.8 e -1.5 , respectivamente. Observamos que, para os três níveis nominais e os diferentes cenários, o intervalo de confiança do tipo assintótico apresentou o menor comprimento médio. Notamos ainda que o intervalo de confiança bootstrap- t apresenta a melhor cobertura, seguido do intervalo ICPE que teve valores muito similares com o do intervalo ICA.

As Figuras 2, 3 e 4 contém histogramas construídos a partir das 10000 estimativas de máxima verossimilhança do parâmetro β_2 , para $n = 40$, $\lambda \approx 128$ e os diferentes cenários para μ_t . As diferentes linhas representam os diferentes intervalos de confiança sob avaliação, sendo seus comprimentos correspondentes aos respectivos comprimentos médios. Os valores abaixo (acima) das linhas são as porcentagens de réplicas em que o verdadeiro valor do parâmetro esteve abaixo (acima) do limite inferior (superior) do intervalo. Esses gráficos foram baseados em (OSPINA; CRIBARI-NETO; VASCONCELLOS, 2006). Através deles é possível verificar que para os diferentes cenários de μ_t , os intervalos analisados são aproximadamente simétricos em torno do verdadeiro valor do parâmetro. Notamos ainda que para $\mu_t \in (0.02, 0.32)$, os intervalos apresentaram-se melhor balanceado quando comparado aos cenários em que $\mu_t \in (0.19, 0.86)$ e $\mu_t \in (0.78, 0.98)$, isto é, as taxas empíricas (% Esquerdo) e (% Direito) coincidem aproximadamente nas caudas da densidade do estimador de máxima verossimilhança de β_2 considerando os diferentes níveis de confiança; isto pode ser verificado também através das Tabelas 7 e 10. De modo geral, o intervalo de confiança bootstrap- t destaca-se como melhor balanceado.

A Tabela 11 apresenta resultados de simulação para o parâmetro de dispersão $\gamma_2 = -2.4$, $n = 40$ e $\lambda = 128$ e os diferentes cenários para μ_t . Notamos que em geral, para os três níveis nominais de confiança e os diferentes cenários considerados, o comprimento médio do

Tabela 7 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para β_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $\mu_t \in (0.02, 0.32)$; $\beta_2 = 1.2$; $n = 40$.

$\lambda \approx 12$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	1.1535	1.6440	0.4905	0.8593	6.61	7.46
	ICBT	1.1270	1.6828	0.5559	0.9016	4.74	5.10
	ICPE	1.1473	1.6443	0.4969	0.8677	5.97	7.26
5%	ICA	1.1066	1.6910	0.5844	0.9190	3.61	4.49
	ICBT	1.0744	1.7445	0.6701	0.9499	2.40	2.61
	ICPE	1.0973	1.6889	0.5916	0.9233	3.20	4.47
1%	ICA	1.0147	1.7828	0.7681	0.9757	0.83	1.60
	ICBT	0.9640	1.8693	0.9053	0.9909	0.34	0.57
	ICPE	1.0025	1.7771	0.7746	0.9761	0.83	1.56
$\lambda \approx 45$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	1.1541	1.6570	0.5029	0.8582	7.23	6.95
	ICBT	1.1272	1.6920	0.5648	0.8962	5.63	4.75
	ICPE	1.1516	1.6575	0.5058	0.8621	7.00	7.30
5%	ICA	1.1059	1.7052	0.5993	0.9187	4.10	4.03
	ICBT	1.0740	1.7551	0.6811	0.9498	2.78	2.24
	ICPE	1.1007	1.7025	0.6018	0.9216	3.77	4.07
1%	ICA	1.0117	1.7993	0.7876	0.9767	1.10	1.23
	ICBT	0.9616	1.8832	0.9216	0.9880	0.70	0.50
	ICPE	1.0038	1.7925	0.7887	0.9766	1.01	1.33
$\lambda \approx 128$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	0.9380	1.4707	0.5327	0.8558	7.01	7.41
	ICBT	0.9072	1.5078	0.6005	0.8987	5.21	4.92
	ICPE	0.9366	1.4743	0.5377	0.8550	7.25	7.25
5%	ICA	0.8869	1.5217	0.6348	0.9166	3.94	4.40
	ICBT	0.8504	1.5750	0.7246	0.9495	2.66	2.39
	ICPE	0.8829	1.5223	0.6394	0.9159	3.98	4.43
1%	ICA	0.7872	1.6214	0.8342	0.9732	1.10	1.58
	ICBT	0.7314	1.7122	0.9807	0.9885	0.61	0.54
	ICPE	0.7814	1.6175	0.8362	0.9733	1.09	1.58

Fonte: próprio autor.

Tabela 8 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para β_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $\mu_t \in (0.02, 0.32)$; $\beta_2 = 1.2$; $n = 80$.

$\lambda \approx 12$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	1.2177	1.5795	0.3618	0.8858	5.39	6.03
	ICBT	1.2100	1.5930	0.3830	0.9011	4.79	5.10
	ICPE	1.2155	1.5789	0.3634	0.8865	5.24	6.11
5%	ICA	1.1831	1.6142	0.4312	0.9375	2.78	3.47
	ICBT	1.1747	1.6334	0.4588	0.9509	2.35	2.56
	ICPE	1.1791	1.6117	0.4326	0.9376	2.61	3.63
1%	ICA	1.1153	1.6819	0.5666	0.9853	0.54	0.93
	ICBT	1.1027	1.7126	0.6100	0.9885	0.53	0.62
	ICPE	1.1101	1.6770	0.5669	0.9838	0.55	1.07
$\lambda \approx 45$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	1.2167	1.5877	0.3710	0.8823	5.98	5.79
	ICBT	1.2076	1.5973	0.3897	0.8952	5.30	5.18
	ICPE	1.2183	1.5884	0.3701	0.8846	5.99	5.55
5%	ICA	1.1812	1.6232	0.4420	0.9346	3.19	3.35
	ICBT	1.1716	1.6388	0.4671	0.9457	2.80	2.63
	ICPE	1.1809	1.6219	0.4409	0.9374	2.97	3.29
1%	ICA	1.1117	1.6927	0.5809	0.9829	0.76	0.95
	ICBT	1.0985	1.7198	0.6213	0.9879	0.62	0.59
	ICPE	1.1102	1.6883	0.5780	0.9812	0.80	1.08
$\lambda \approx 128$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	1.0065	1.4012	0.3946	0.8838	5.93	5.69
	ICBT	0.9929	1.4098	0.4169	0.9034	4.92	4.74
	ICPE	1.0099	1.4062	0.3963	0.8815	6.53	5.32
5%	ICA	0.9687	1.4390	0.4703	0.9349	3.29	3.22
	ICBT	0.9547	1.4543	0.4996	0.9523	2.60	2.17
	ICPE	0.9701	1.4421	0.4720	0.9336	3.49	3.15
1%	ICA	0.8948	1.5128	0.6180	0.9827	0.79	0.94
	ICBT	0.8764	1.5416	0.6651	0.9883	0.63	0.54
	ICPE	0.8947	1.5130	0.6183	0.9808	0.90	1.02

Fonte: próprio autor.

Tabela 9 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para β_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $\mu_t \in (0.02, 0.32)$; $\beta_2 = 1.2$; $n = 120$.

$\lambda \approx 12$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	1.2494	1.5492	0.2998	0.8899	5.21	5.80
	ICBT	1.2457	1.5568	0.3111	0.8997	4.90	5.13
	ICPE	1.2481	1.5486	0.3004	0.8892	5.18	5.90
5%	ICA	1.2206	1.5779	0.3573	0.9414	2.65	3.21
	ICBT	1.2172	1.5891	0.3719	0.9478	2.55	2.67
	ICPE	1.2180	1.5757	0.3577	0.9404	2.55	3.41
1%	ICA	1.1645	1.6341	0.4695	0.9863	0.54	0.83
	ICBT	1.1598	1.6516	0.4918	0.9887	0.51	0.62
	ICPE	1.1611	1.6299	0.4688	0.9859	0.42	0.99
$\lambda \approx 45$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	1.2475	1.5544	0.3069	0.8899	5.30	5.71
	ICBT	1.2425	1.5588	0.3163	0.8956	4.89	5.55
	ICPE	1.2492	1.5551	0.3058	0.8921	5.19	5.60
5%	ICA	1.2181	1.5838	0.3657	0.9422	2.67	3.11
	ICBT	1.2135	1.5917	0.3781	0.9463	2.64	2.73
	ICPE	1.2186	1.5827	0.3642	0.9418	2.64	3.18
1%	ICA	1.1606	1.6412	0.4806	0.9865	0.56	0.79
	ICBT	1.1548	1.6554	0.5006	0.9882	0.62	0.56
	ICPE	1.1604	1.6379	0.4775	0.9848	0.58	0.94
$\lambda \approx 128$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	1.0442	1.3705	0.3263	0.8877	5.86	5.37
	ICBT	1.0347	1.3722	0.3375	0.8998	5.13	4.89
	ICPE	1.0490	1.3758	0.3268	0.8854	6.63	4.83
5%	ICA	1.0130	1.4018	0.3888	0.9402	3.19	2.79
	ICBT	1.0040	1.4077	0.4036	0.9513	2.77	2.10
	ICPE	1.0161	1.4054	0.3893	0.9392	3.57	2.51
1%	ICA	0.9519	1.4629	0.5110	0.9864	0.64	0.72
	ICBT	0.9418	1.4762	0.5344	0.9884	0.65	0.51
	ICPE	0.9540	1.4641	0.5101	0.9845	0.82	0.73

Fonte: próprio autor.

intervalo de confiança do tipo assintótico é menor do que os comprimentos médios dos demais intervalos, seguido do comprimento médio do intervalo de confiança bootstrap- t . O intervalo de confiança bootstrap percentil apresenta a melhor cobertura e os valores do intervalo ICBT são similares ao do intervalo ICPE.

Para os diferentes cenários de μ_t , as Figuras 5, 6 e 7 mostram que apenas o intervalo de confiança assintótico é aproximadamente simétrico em torno do verdadeiro valor do parâmetro. O intervalo de confiança bootstrap- t é levemente assimétrico ao redor de γ_2 . Em relação ao intervalo de confiança bootstrap percentil, este apresenta assimetria muito forte, principalmente para o nível nominal 99%. Observamos ainda que o intervalo de confiança assintótico apresenta forte não balanceamento, dado que as taxas % Direito são acentuadamente superiores às taxas observadas % Esquerdo. Porém, os intervalos de confiança bootstrap- t e bootstrap percentil são aproximadamente balanceados para todos os níveis nominais e cenários.

Portanto, com base nos resultados apresentados sugere-se o uso do intervalo de confiança bootstrap- t que apresentou melhor cobertura e balanceamento.

Tabela 10 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap- t (ICBT) e percentil (ICPE) para β_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $n = 40$; $\lambda \approx 128$.

$\mu_t \in (0.19, 0.86), \beta_2 = -1.8.$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	-1.9943	-1.6095	0.3849	0.8615	5.02	8.83
	ICBT	-1.9982	-1.5596	0.4386	0.9063	4.57	4.80
	ICPE	-1.9949	-1.6028	0.3921	0.8655	5.15	8.30
5%	ICA	-2.0312	-1.5726	0.4586	0.9207	2.46	5.47
	ICBT	-2.0350	-1.5051	0.5299	0.9515	2.27	2.58
	ICPE	-2.0303	-1.5628	0.4674	0.9234	2.58	5.08
1%	ICA	-2.1033	-1.5005	0.6027	0.9710	0.49	2.41
	ICBT	-2.1107	-1.3925	0.7182	0.9889	0.48	0.63
	ICPE	-2.0952	-1.4805	0.6148	0.9732	0.65	2.03
$\mu_t \in (0.78, 0.98), \beta_2 = -1.5.$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	-1.6143	-1.3839	0.2304	0.8587	6.13	8.00
	ICBT	-1.6244	-1.3636	0.2608	0.9008	4.60	5.32
	ICPE	-1.6127	-1.3790	0.2337	0.8624	6.39	7.37
5%	ICA	-1.6364	-1.3619	0.2745	0.9173	3.27	5.00
	ICBT	-1.6482	-1.3340	0.3142	0.9475	2.47	2.78
	ICPE	-1.6344	-1.3558	0.2786	0.9178	3.68	4.54
1%	ICA	-1.6795	-1.3187	0.3607	0.9717	1.01	1.82
	ICBT	-1.6979	-1.2738	0.4240	0.9875	0.60	0.65
	ICPE	-1.6747	-1.3086	0.3660	0.9726	1.19	1.55

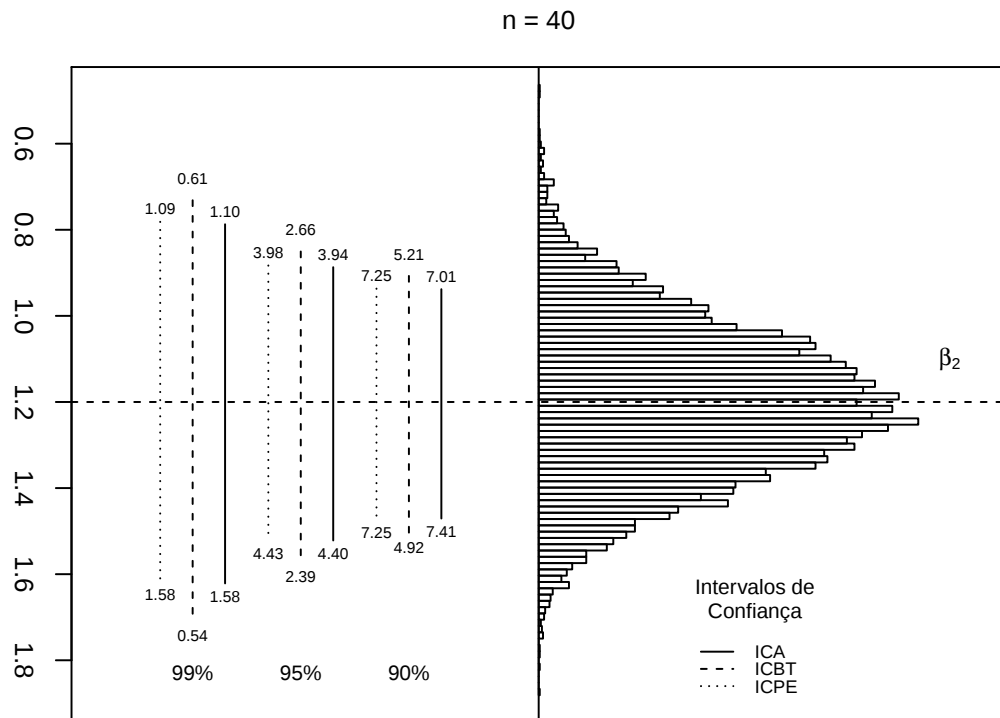
Fonte: próprio autor.

Tabela 11 – Limites inferiores e superiores, tamanho, probabilidades empíricas de cobertura (Cobertura) e percentuais de não cobertura inferior (% Esquerdo) e superior (% Direito) dos estimadores intervalares assintótico (ICA), bootstrap-t (ICBT) e percentil (ICPE) para γ_2 no modelo: $\text{logit}(\mu_t/1 - \mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4}$ e $\log(\sigma_t^2) = \gamma_1 + z_t^{\gamma_2}$, $t = 1, \dots, n$; $n = 40$; $\lambda \approx 128$.

$\mu_t \in (0.02, 0.32), \quad \gamma_2 = -2.4.$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	-2.7589	-2.0444	0.7144	0.8632	4.52	9.16
	ICBT	-2.7488	-1.9600	0.7888	0.8925	5.40	5.35
	ICPE	-2.7558	-1.9421	0.8138	0.8970	5.14	5.16
5%	ICA	-2.8273	-1.9760	0.8513	0.9227	2.07	5.66
	ICBT	-2.8104	-1.8717	0.9388	0.9396	3.25	2.79
	ICPE	-2.8119	-1.8197	0.9923	0.9482	3.02	2.16
1%	ICA	-2.9611	-1.8422	1.1188	0.9791	0.22	1.87
	ICBT	-2.9214	-1.6958	1.2256	0.9804	1.22	0.74
	ICPE	-2.9137	-1.5208	1.3930	0.9875	0.99	0.26
$\mu \in (0.19, 0.86), \quad \gamma_2 = -2.4.$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	-2.7653	-2.0544	0.7109	0.8609	4.04	9.87
	ICBT	-2.7506	-1.9569	0.7937	0.8967	5.03	5.30
	ICPE	-2.7713	-1.9608	0.8106	0.8983	4.27	5.90
5%	ICA	-2.8334	-1.9863	0.8471	0.9205	1.89	6.06
	ICBT	-2.8138	-1.8677	0.9461	0.9434	2.85	2.81
	ICPE	-2.8274	-1.8365	0.9909	0.9484	2.46	2.70
1%	ICA	-2.9665	-1.8532	1.1133	0.9772	0.20	2.08
	ICBT	-2.9262	-1.6884	1.2378	0.9830	0.95	0.75
	ICPE	-2.9302	-1.5360	1.3942	0.9879	0.84	0.37
$\mu_t \in (0.78, 0.98), \quad \gamma_2 = -2.4.$							
α	Intervalos	Inferior	Superior	Tamanho	Cobertura	% Esquerdo	% Direito
10%	ICA	-2.7655	-2.0547	0.7108	0.8589	4.08	10.03
	ICBT	-2.7507	-1.9578	0.7929	0.8987	4.81	5.32
	ICPE	-2.7712	-1.9609	0.8103	0.8936	4.23	6.41
5%	ICA	-2.8336	-1.9866	0.8470	0.9173	2.07	6.20
	ICBT	-2.8139	-1.8691	0.9447	0.9454	2.69	2.77
	ICPE	-2.8270	-1.8365	0.9905	0.9476	2.56	2.68
1%	ICA	-2.9667	-1.8535	1.1131	0.9765	0.27	2.08
	ICBT	-2.9282	-1.6928	1.2354	0.9824	1.03	0.73
	ICPE	-2.9281	-1.5294	1.3987	0.9868	0.94	0.38

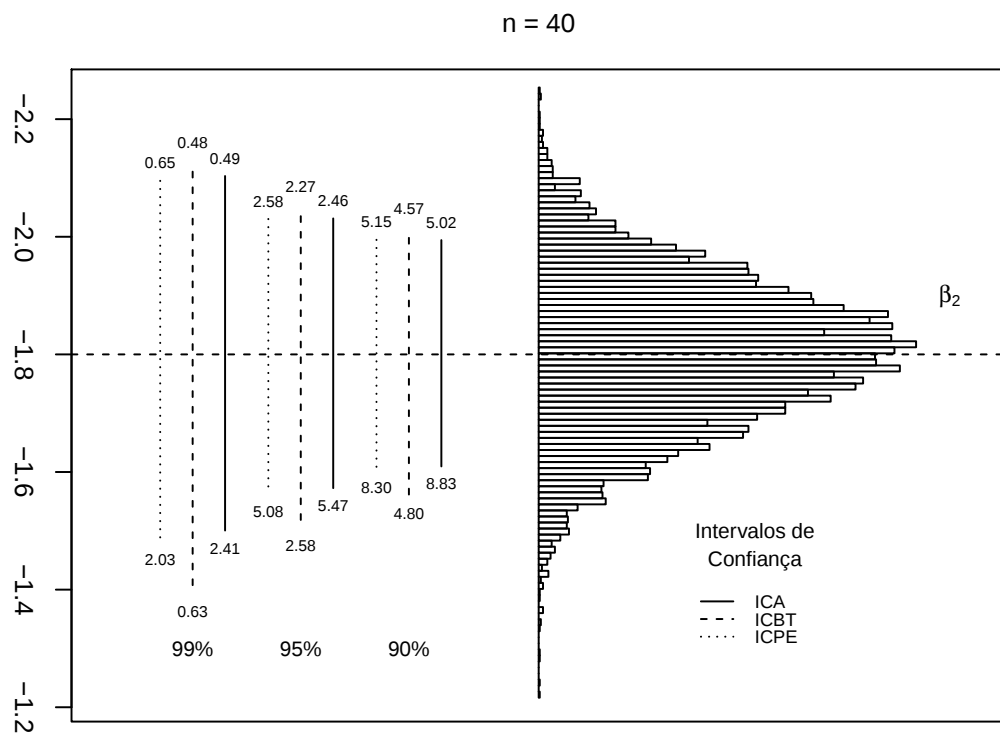
Fonte: próprio autor.

Figura 2 – Estimação intervalar para β_2 com $n = 40$, $\mu_t \in (0.02, 0.32)$ e $\lambda \approx 128$.



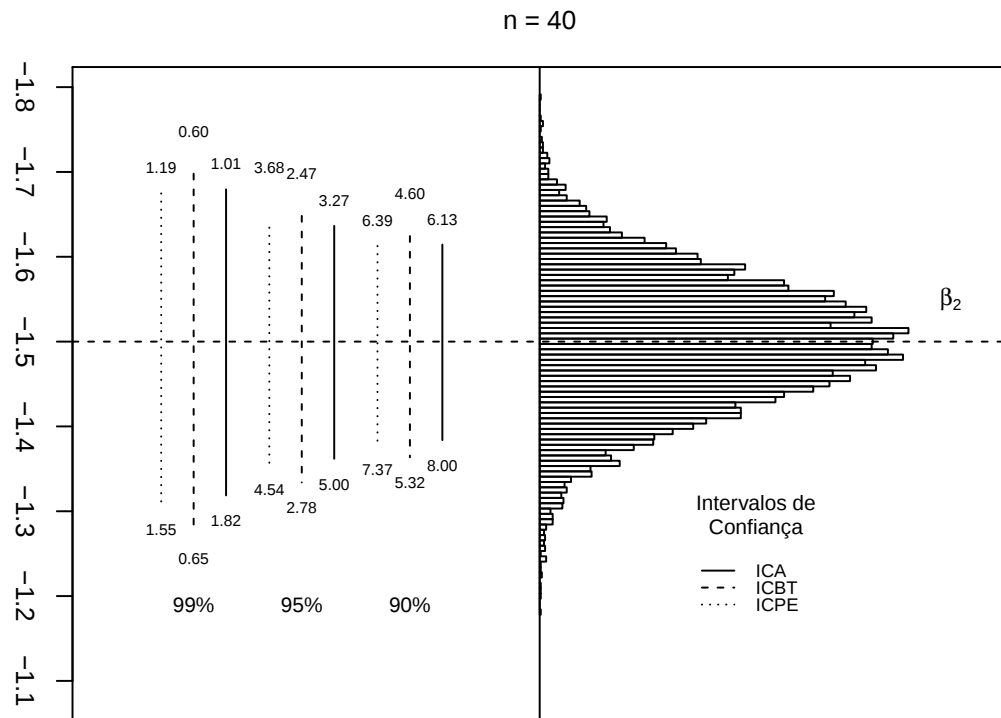
Fonte: próprio autor.

Figura 3 – Estimação intervalar para β_2 com $n = 40$, $\mu_t \in (0.19, 0.86)$ e $\lambda \approx 128$.



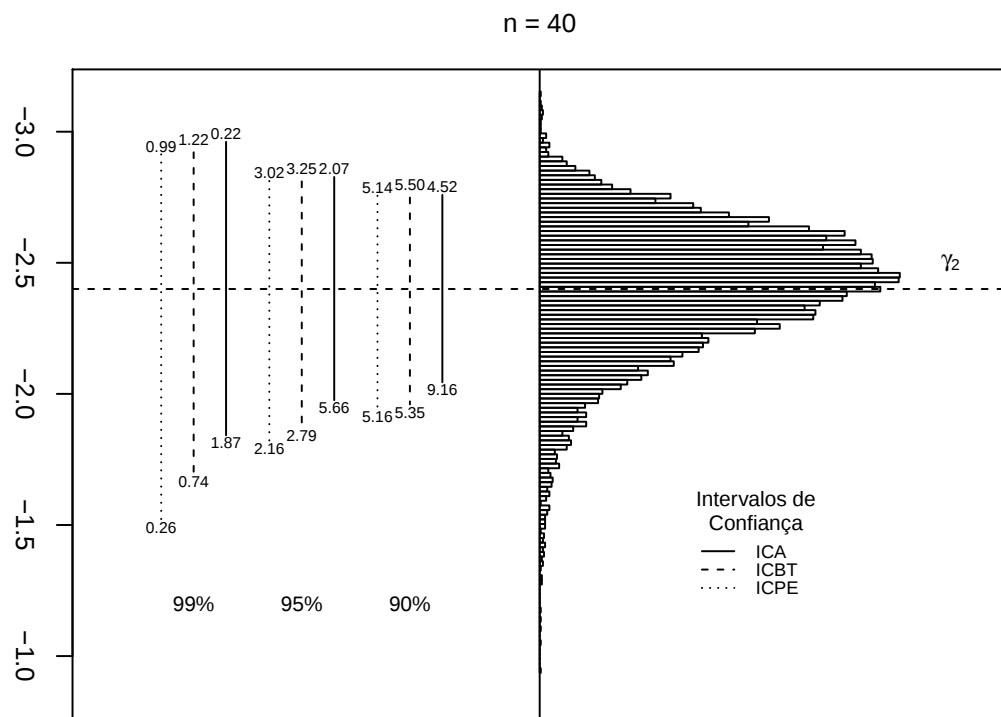
Fonte: próprio autor.

Figura 4 – Estimação intervalar para β_2 com $n = 40$, $\mu_t \in (0.78, 0.98)$ e $\lambda \approx 128$.



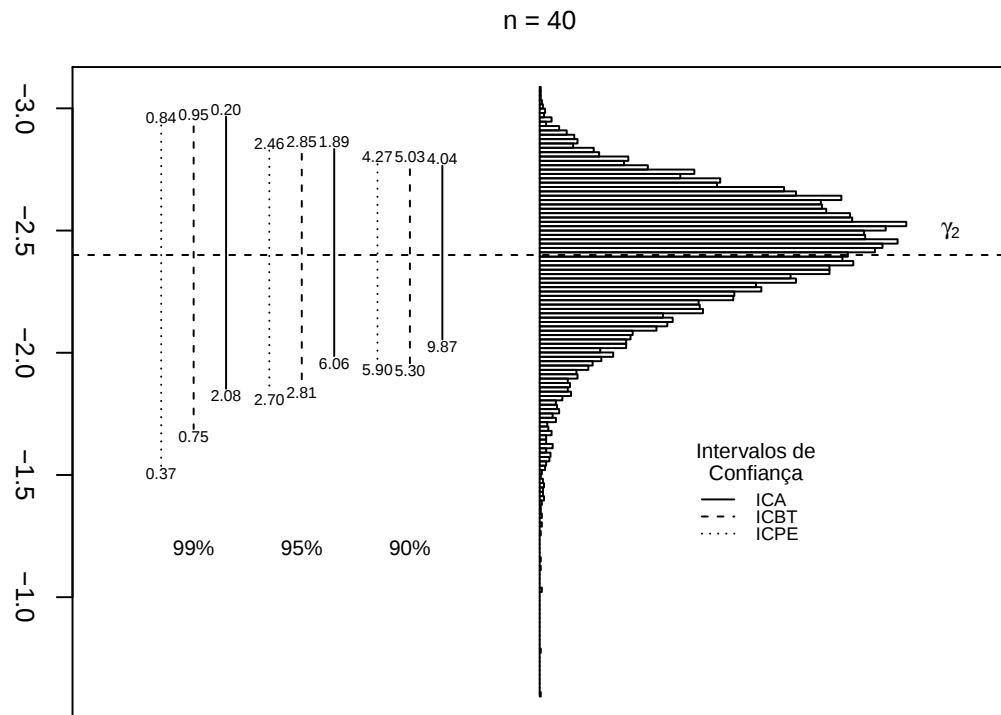
Fonte: próprio autor.

Figura 5 – Estimação intervalar para γ_2 com $n = 40$, $\mu_t \in (0.02, 0.32)$ e $\lambda \approx 128$.



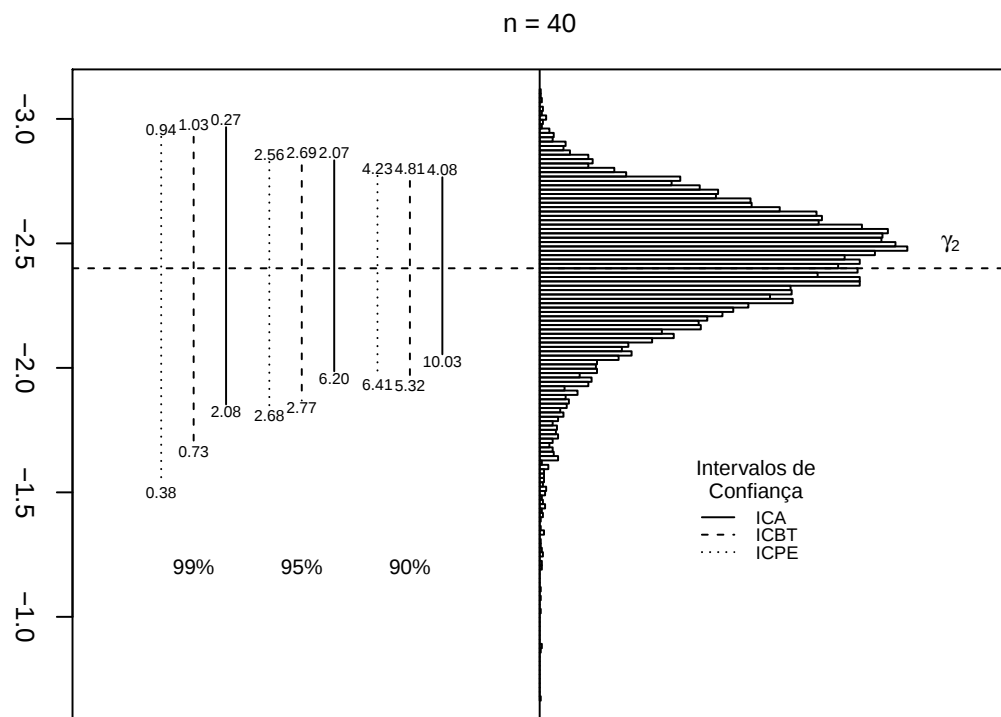
Fonte: próprio autor.

Figura 6 – Estimação intervalar para γ_2 com $n = 40$, $\mu_t \in (0.19, 0.86)$ e $\lambda \approx 128$.



Fonte: próprio autor.

Figura 7 – Estimação intervalar para γ_2 com $n = 40$, $\mu_t \in (0.78, 0.98)$ e $\lambda \approx 128$.



Fonte: próprio autor.

3.7 APLICAÇÃO: DADOS DE CRAQUEAMENTO CATALÍTICO FLUIDO (FCC)

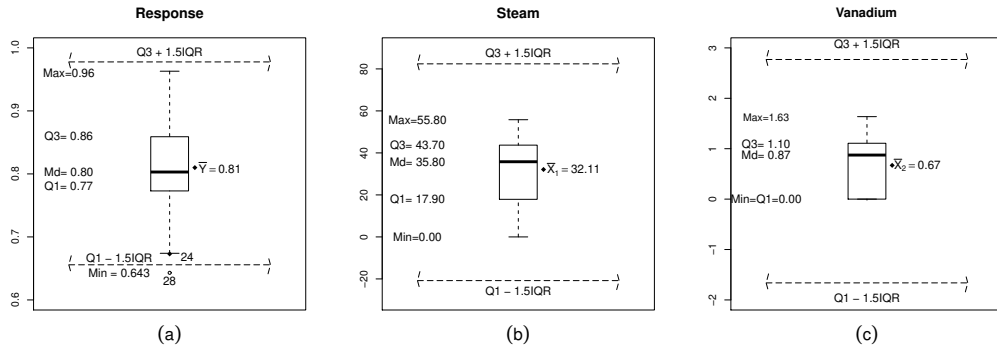
Os dados desta aplicação são do Departamento de Química da Universidade Nacional da Colômbia (SALAZAR, 2005). O processo craqueamento catalítico fluido conhecido como Fluid Catalytic Cracking (FCC) é usado para converter hidrocarbonetos de alto peso molecular em pequenas moléculas de maior valor comercial através do contato destes com um catalisador. Este processo é muitas vezes considerado o coração de uma refinaria, uma vez que permite que a produção seja adaptada aos produtos de maior demanda e/ou alta rentabilidade (SALAZAR, 2005). O catalisador do processo consiste em partículas finas de 10 a 150 microns, facilmente fluidizável tendo como componente principal o zeólito Y (SALAZAR, 2005). Outra substância importante que participa do processo de catalização é o vanádio. Este componente químico é conhecido por participar da destruição do catalisador, reduzindo a superfície ativa, a seletividade e a cristalinidade do zeólito Y especialmente na presença de vapor. Sabe-se que cada 1000 ppm (partes por milhão) de vanádio no catalisador reduz a produção de gasolina em cerca de 2.3%. O processo também depende da temperatura, que deve ficar próxima de $720^{\circ}C$ (SALAZAR, 2005).

Esses dados foram modelados pela primeira vez por Espinheira e Silva (2019). O objetivo é modelar a variável resposta porcentagem de cristalinidade do zeólito Y (y_t) através das covariáveis vapor d'água (x_2), temperatura do processo (x_3) e concentração de vanádio (x_4). O conjunto de dados é composto por 28 observações. Inicialmente, Espinheira e Silva (2019) realizaram uma análise descritiva da variável resposta e covariáveis a qual vamos apresentar na Figura 8. Nesta figura apresentamos algumas estatísticas descritivas da variável resposta (Response) e das covariáveis vapor d'água (Steam) e concentração de vanádio (Vanadium) usando boxplots.

Com base na Figura 8(a) nota-se que a variável resposta está concentrada perto da parte superior do intervalo unitário. Adicionalmente, observa-se a presença de um outlier, a observação 28, a qual assume o menor valor da resposta igual a 0.643. A Figura 8(b) revela a alta dispersão da covariável vapor d'água (x_2) que pode assumir valores entre -20 e 80 . Na Figura 8(c) vemos que o mínimo da covariável vanádio é igual ao primeiro quantil assumindo valor zero, ou seja, um quarto dos valores assumidos pela covariável é igual a zero. A covariável temperatura assume dois valores, 700° ou 760° , e cada um desses valores representa 50% dos dados.

Com base na Figura 8(b) Espinheira e Silva (2019) argumentam que parece interessante

Figura 8 – Boxplots dos valores da resposta (a) e das covariáveis vapor d'água (b) e vanádio (c). Dados de craqueamento catalítico fluido (FCC).



Fonte: próprio autor.

usar uma função para estabilizar essa dispersão da covariável. Uma função plausível é $\log(x_2 + \beta)$ devido à presença de zeros. Outras funções plausíveis são $1/(x_2 + \beta)$ ou $x_2/(x_2 + \beta)$. Todas elas são funções não lineares nos parâmetros. Dessa forma, os autores escolhem o modelo simplex definido por

$$\Phi^{-1}(\mu_t) = \beta_1 + \beta_2 \frac{x_{t2}}{x_{t2} + \beta_3} + \beta_4 x_{t3} + \beta_5 \sqrt{x_{t4}} \quad \text{e} \quad \log(\sigma_t^2) = \gamma_1 + \gamma_2 x_{t4}^2, \quad (3.4)$$

com $t = 1, \dots, 28$. Para mais informações sobre a construção desse modelo, ver Espinheira e Silva (2019).

Notamos que em (3.4) há mais parâmetros que covariáveis. Para obter $(\beta_1^{(0)}, \beta_2^{(0)}, \beta_4^{(0)}, \beta_5^{(0)})^\top$, Espinheira e Silva (2019) sugerem usar $(X^\top X)^{-1} X^\top g(y)$ em que a t -linha da matriz X é definida como $x_t^\top = (1, x_{t2}/(x_{t2} + \beta_3^{(0)}), x_{t3}, \sqrt{x_{t4}})$, $t = 1, \dots, 28$. No entanto, é necessário fornecer um valor numérico para $\beta_3^{(0)}$. Baseado na Figura 8(b), os autores consideram as equações $[x_2 + \beta_3^{(0)} = -20]$ e $[x_2 + \beta_3^{(0)} = 80]$, $x_2 = 0$ e $x_2 = 55.8$, o que implica que $\beta_3^{(0)} \in (-20, 80)$ e $\beta_3^{(0)} \in (-75.80, 24.2)$, respectivamente. Assim, é sugerido $\beta_3^{(0)} = -20$ o qual conduziu a estimativas de parâmetros plausíveis com erros-padrão finitos. Portanto,

$$\beta_{NL}^{(0)} = (\tilde{\mathcal{X}}^{(0)\top} \tilde{\mathcal{X}}^{(0)})^{-1} \tilde{\mathcal{X}}^{(0)\top} [g(y) - f_1(x_t^\top, \beta_L^{(0)})],$$

em que $f_1(x_t^\top, \beta_L^{(0)}) = \beta_1^{(0)} + \beta_2^{(0)} x_{t2}/(x_{t2} - 20) + \beta_4^{(0)} x_{t3} + \beta_5^{(0)} \sqrt{x_{t4}}$ e a t -ésima linha da matriz $\tilde{\mathcal{X}}^{(0)} = [\partial \eta / \partial \beta]_{\beta = \beta^{(0)}}$ é definida como $(1, x_{t2}/(x_{t2} - 20), -\beta_2^{(0)} x_{t2}/((x_{t2} - 20)^{-2}), x_{t3}, \sqrt{x_{t4}})$.

As estimativas dos parâmetros foram obtidas através da maximização da função de log-verossimilhança com base no método de otimização não linear Escore de Fisher. Na Tabela 12 apresentamos as estimativas de máxima verossimilhança, as estimativas com viés corrigido

via bootstrap, seus respectivos erros-padrão (EP) e os p -valores associados aos testes Z para a significância dos parâmetros do modelo.

Inicialmente notamos que todos os parâmetros foram significativos estatisticamente ao nível de 5%, tanto para o submodelo da média quanto para o submodelo da dispersão. Notamos que as estimativas de máxima verossimilhança para os parâmetros do submodelo da média não diferem muito das estimativas corrigidas via esquema bootstrap. Porém, as estimativas corrigidas apresentam erros-padrão inferiores aos de MV com exceção apenas para o parâmetro β_1 . Em relação aos parâmetros do submodelo da dispersão, percebemos que as estimativas de máxima verossimilhança e sua versão corrigida apresentam valores diferentes porém erros-padrão iguais.

Tabela 12 – Estimativas de máxima verossimilhança $\hat{\theta}$, estimativas com viés corrigido via bootstrap ($\bar{\theta}$), erros-padrão (EP) e p -valores associados aos testes Z para os parâmetros do Modelo 3.4. Dados de craqueamento catalítico fluido (FCC).

Parâmetros	β_1	β_2	β_3	β_4	β_5	γ_1	γ_2
$\hat{\theta}$	1.3994	-0.0609	-27.8429	-0.1820	-0.4201	0.7561	-1.1502
$\bar{\theta}$	1.4109	-0.0688	-27.7823	-0.1802	-0.4276	0.8981	-0.9443
$EP(\hat{\theta})$	0.0769	0.0227	3.8788	0.0629	0.0747	0.3698	0.3143
$EP(\bar{\theta})$	0.0816	0.0185	3.1853	0.0441	0.0661	0.3698	0.3143
P -valor	< 0.0001	0.0074	< 0.0001	0.0038	< 0.0001	0.0409	0.0003

Fonte: próprio autor.

Na Tabela 13 encontram-se as estimativas intervalares dos parâmetros do modelo com níveis nominais iguais a 90%, 95% e 99%. Observamos que não há grande diferença nos limites inferiores e superiores dos intervalos de confiança para os parâmetros β_1 , β_2 , β_4 e β_5 . Em relação ao parâmetro β_3 , percebemos que o intervalo de confiança bootstrap- t apresenta comportamento diferente dos outros dois intervalos. Notamos ainda que para os níveis nominais 95%, 99% e parâmetro β_4 , o intervalo de confiança bootstrap- t contém o valor zero, o que não acontece com o nível nominal 90% cujo intervalo é dado por $(-0.3462; -0.0029)$. Observamos que este mesmo fato acontece para o nível nominal 99%, parâmetro β_2 e intervalos de confiança bootstrap- t e percentil. Essa informação em relação a β_2 é importante pois a possibilidade de $\beta_2 = 0$ implica a exclusão do termo $(x_2/x_2 + \beta_3)$, ou seja, implica tanto em x_2 ser considerada não importante para a explicação da resposta quanto em o modelo ser linear. Neste ponto, relembramos como ao nível de 99% os resultados de simulação apresen-

tarem a possibilidade dessa conclusão ser enganosa e, portanto, a cautela sugere que sejam considerados os resultados aos níveis de 90% e 95%.

Em relação aos intervalos de confiança para os parâmetros do submodelo da dispersão, notamos que os valores são semelhantes tanto para os limites inferiores quanto para os superiores. No entanto, percebemos que alguns intervalos de confiança contêm o valor zero. Por exemplo, para o parâmetro γ_1 podemos encontrar nos seguintes casos: nível nominal 90% e intervalo bootstrap percentil; nível nominal 95% e intervalos bootstrap- t e percentil; nível nominal 99% e em todos os intervalos. E para o parâmetro γ_2 : nível nominal 95% e intervalo bootstrap- t ; nível nominal 99% e intervalos bootstrap- t e percentil.

Tabela 13 – Estimativas intervalares assintótica (ICA), bootstrap-t (ICBT) e percentil (ICPE) para os parâmetros do Modelo 3.4. Dados de craqueamento catalítico fluido (FCC)

90%			
Parâmetros	ICA	ICBT	ICPE
β_1	(1.2729; 1.5259)	(1.2352; 1.5473)	(1.2501; 1.5313)
β_2	(-0.0983; -0.0235)	(-0.1170; -0.0203)	(-0.0858; -0.0112)
β_3	(-34.2230; -21.4628)	(-38.5221; -15.8487)	(-34.2532; -21.1564)
β_4	(-0.2854; -0.0786)	(-0.3462; -0.0029)	(-0.3081; -0.0689)
β_5	(-0.5430; -0.2971)	(-0.5541; -0.2680)	(-0.5442; -0.2838)
γ_1	(0.1479; 1.3644)	(0.1107; 1.7900)	(-0.2778; 1.4015)
γ_2	(-1.6671; -0.6333)	(-1.7556; -0.1909)	(-2.1096; -0.5449)
95%			
Parâmetros	ICA	ICBT	ICPE
β_1	(1.2487; 1.5501)	(1.1866; 1.5800)	(1.1857; 1.5667)
β_2	(-0.1055; -0.0164)	(-0.1464; -0.0085)	(-0.0945; -0.0044)
β_3	(-35.4453; -20.2405)	(-40.9290; -9.5443)	(-34.7452; -20.4744)
β_4	(-0.3052; -0.0588)	(-0.3817; 0.0610)	(-0.3311; -0.0383)
β_5	(-0.5666; -0.2736)	(-0.5785; -0.2181)	(-0.5663; -0.2607)
γ_1	(0.0313; 1.4809)	(-0.0420; 1.9919)	(-0.4796; 1.5542)
γ_2	(-1.7662; -0.5343)	(-1.9857; 0.0006)	(-2.3011; -0.3148)
99%			
Parâmetros	ICA	ICBT	ICPE
β_1	(1.2013; 1.5974)	(1.0927; 1.6128)	(1.1423; 1.6074)
β_2	(-0.1195; -0.0024)	(-0.2258; 0.0183)	(-0.1042; 0.0098)
β_3	(-37.8341; -17.8517)	(-46.4448; -4.6130)	(-37.0234; -18.3824)
β_4	(-0.3439; -0.0201)	(-0.4656; 0.1477)	(-0.3903; -0.0132)
β_5	(-0.6126; -0.2275)	(-0.6602; -0.1735)	(-0.6198; -0.2191)
γ_1	(-0.1964; 1.7087)	(-0.2436; 2.2129)	(-0.7007; 1.7558)
γ_2	(-1.9597; -0.3408)	(-2.3156; 0.3384)	(-2.6389; 0.0151)

Fonte: próprio autor.

4 TESTES DE HIPÓTESES

4.1 INTRODUÇÃO

Outro aspecto importante na inferência estatística são os testes de hipóteses, que têm por objetivo tomar decisões sobre os parâmetros do modelo com uma probabilidade de erro pré-estabelecida. Inferências em grandes amostras no modelo de regressão simplex não linear (ESPINHEIRA; SILVA, 2019) podem ser realizadas através dos testes da razão de verossimilhanças generalizada (WILKS, 1938), escore (RAO, 1948), Wald (WALD, 1943) e gradiente (TERRELL, 2002). As estatísticas destes testes baseiam-se na função de verossimilhança, que apresenta várias propriedades de otimalidade. Sob condições usuais de regularidade (ver Apêndice B) e sob a hipótese nula, as distribuições limite destas estatísticas são qui-quadrado, ou seja, as estatísticas são assintoticamente equivalentes.

Estes testes são realizados a partir de valores críticos válidos para grandes amostras, o que pode causar considerável distorção no tamanho do teste em amostras pequenas ou moderadas. Esta distorção pode ser contornada através de correções entre as quais baseadas na técnica de reamostragem bootstrap. Na literatura, existem vários testes utilizando este método. Aqui, faremos uso do teste da razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (ROCKE, 1989). Adicionalmente, ainda com base na estatística RV vamos utilizar o bootstrap duplo rápido (DAVIDSON; MACKINNON, 2000) para a obtenção do p -valor empírico do teste, o que tipicamente reduz a distorção de tamanho associada ao teste assintótico usual.

O objetivo deste capítulo é corrigir o teste da razão de verossimilhanças generalizada por meio da correção de Bartlett bootstrap e bootstrap duplo rápido em pequenas amostras no modelo de regressão simplex não linear. Além disso, avaliar numericamente via simulações de Monte Carlo as correções propostas e compará-las com a versão não corrigida e com os testes escore, Wald e gradiente. Por último, aplicar os testes avaliados ao conjunto de dados modelados por Espinheira e Silva (2019). Foi considerado testes sobre os parâmetros do submodelo da média e sobre os parâmetros que indexam o submodelo da dispersão.

4.2 TESTES ASSINTÓTICOS

Quando a distribuição exata da estatística de um teste de hipótese é analiticamente difícil ou até impossível de ser definida uma alternativa é utilizar distribuições aproximadas o que

pode ser feito a partir da teoria assintótica. Os testes assintóticos mais conhecidos na literatura, chamados também de testes clássicos, são os da razão de verossimilhanças generalizada, escore e Wald (BUSE, 1982). Adicionalmente, o teste gradiente proposto por Terrell (2002) é um teste novo comparado com os demais, mas que vem ganhando bastante atenção nos últimos anos.

Considere θ_p um subconjunto de θ e o teste da hipótese nula

$$H_0 : \theta_p = \theta_p^{(0)} \quad \text{versus a hipótese alternativa} \quad H_1 : \theta_p \neq \theta_p^{(0)}. \quad (4.1)$$

Definindo $\theta_{(H_0)} = \theta^{(0)}$ a restrição imposta por H_0 , temos que $\theta^{(0)} \in \Theta_0$ e Θ_0 é o espaço paramétrico restrito a hipótese nula, $\Theta_0 \subset \Theta$. Adicionalmente, vamos denotar $\tilde{\theta}$, $\hat{\theta}$, $\tilde{K}_{pp}^{-1}$ e $\hat{K}_{pp}$ como o estimador de máxima verossimilhança restrito a H_0 , o estimador de máxima verossimilhança irrestrito, a inversa da matriz $p \times p$ de informação de Fisher avaliada em $\tilde{\theta}$ e a inversa da matriz $p \times p$ de informação de Fisher avaliada em $\hat{\theta}$, respectivamente. A seguir vamos descrever os testes que vamos utilizar para testar (4.1).

4.2.1 Teste da Razão de Verossimilhanças Generalizada

A função de verossimilhança contém toda a informação relevante para fazer inferência acerca de um vetor de parâmetros de interesse. O teste da razão de verossimilhanças generalizada (WILKS, 1938) se baseia na comparação dos ajustes de dois modelos, o modelo nulo, que é obtido impondo-se a restrição em teste e, o modelo irrestrito, que é o modelo sem restrições.

A estatística do teste da razão de verossimilhanças generalizada para testar (4.1) é definida como

$$\Lambda = \frac{L(\tilde{\theta}; y)}{L(\hat{\theta}; y)},$$

em que $L(\tilde{\theta}; y)$ é a função de verossimilhança avaliada na amostra e em $\tilde{\theta}$ e $L(\hat{\theta}; y)$ a função de verossimilhança avaliada na amostra e em $\hat{\theta}$. A estatística Λ revela de qual modelo é mais verossímil que a amostra tenha vindo. A razão de verossimilhanças (ou seu logaritmo) pode ser utilizada para decidir se a hipótese nula em teste será ou não rejeitada.

A distribuição nula exata de Λ é tipicamente difícil de ser obtida. Por isso, utiliza-se uma transformação conveniente de Λ dada por

$$RV = -2 \log \Lambda = 2\{l(\hat{\theta}) - l(\tilde{\theta})\},$$

em que $l(\hat{\theta})$ e $l(\tilde{\theta})$ são os logaritmos da função de verossimilhança avaliados em $\hat{\theta}$ e $\tilde{\theta}$, estimadores de máxima verossimilhança irrestrito e restrito a H_0 , respectivamente.

4.2.2 Teste Escore de Rao

Outro teste presente na literatura para medir a distância entre as hipóteses nula e alternativa é o teste escore proposto por Rao (1948). Este teste requer somente estimação sob a hipótese nula, o que o torna mais atraente quando a estimação sob a hipótese alternativa exige um extenso trabalho. Sua estatística é dada por

$$S_R = \tilde{U}_p^\top \tilde{K}_{pp}^{-1} \tilde{U}_p,$$

em que $\tilde{U}_p$ é o vetor escore avaliado em $\tilde{\theta}$ e $\tilde{K}_{pp}^{-1}$ é a inversa da matriz $p \times p$ de informação de Fisher avaliada também em $\tilde{\theta}$.

4.2.3 Teste de Wald

O terceiro teste que vamos apresentar é o mais tipicamente utilizado em associação com a estimação pontual por máxima verossimilhança denominado de teste de Wald (WALD, 1943). A estatística do teste para testar a hipótese (4.1) é definida como

$$W = (\hat{\theta}_p - \theta_p^{(0)})^\top \hat{K}_{pp} (\hat{\theta}_p - \theta_p^{(0)}),$$

em que $\hat{\theta}_p$ é o estimador de máxima verossimilhança de θ_p e $\hat{K}_{pp}$ é a matriz $p \times p$ de informação de Fisher avaliada no estimador de máxima verossimilhança irrestrito $\hat{\theta}$.

4.2.4 Teste Gradiente

Como uma alternativa aos testes clássicos, Terrell (2002) propôs uma nova estatística, chamada de gradiente. Esta estatística foi deduzida a partir das estatísticas de Wald e escore de Rao (VARGAS; FERRARI; LEMONTE, 2013), e comparada com as estatísticas clássicas, é muito atrativa pela simplicidade de seu cálculo. Sua obtenção não envolve nenhum cálculo matricial como produto e inversa de matrizes, que em problemas complexos podem representar altos custos computacionais (VARGAS; FERRARI; LEMONTE, 2013).

A estatística gradiente S_G para testar a hipótese (4.1) é dada da seguinte forma:

$$S_G = \tilde{U}_p^\top (\hat{\theta}_p - \theta_p^{(0)}),$$

em que $\hat{\theta}_p$ é o estimador de máxima verossimilhança de θ_p e $\tilde{U}_p$ é o vetor escore avaliado em $\tilde{\theta}$. Para mais informações sobre esta estatística e suas propriedades, ver Terrell (2002) e Lemonte

(2016).

Sob condições usuais de regularidade (ver Apêndice B), temos que

$$RV, S_R, W, S_G \xrightarrow{D} \chi_p^2,$$

sendo p a dimensão do vetor de parâmetros em H_0 . Assim, os quatro testes podem ser realizados usando valores críticos aproximados da distribuição χ_p^2 . Sendo α o nível de significância, os testes que rejeitarão H_0 em (4.1) serão aqueles cujos valores observados de suas estatísticas forem maiores que o quantil $1 - \alpha$ da distribuição qui-quadrado com p graus de liberdade.

Existem algumas semelhanças entre as estatísticas de testes apresentadas, por exemplo, a utilização da função de verossimilhança e a mesma distribuição assintótica sob condições usuais de regularidade. Note que o teste escore é o único que requer estimação somente sob a hipótese nula, e que a estimação para a realização do teste de Wald baseia-se apenas no modelo irrestrito. O teste de Wald foi construído para competir com o teste da razão de verossimilhanças generalizada, no entanto, os pesquisadores notaram que o mesmo tende a rejeitar H_0 muito mais do que deveria. Muggeo e Lovison (2014) apresentam uma interpretação geométrica dos testes da razão de verossimilhanças generalizada, escore, Wald e gradiente, mostrando suas respectivas conexões.

4.3 TESTES DE HIPÓTESES EM REGRESSÃO SIMPLEX NÃO LINEAR SOBRE O VETOR DE PARÂMETROS β

Seja $y = (y_1, y_2, \dots, y_n)^\top$ um vetor n -dimensional de variáveis aleatórias independentes, em que cada y_t , $t = 1, \dots, n$, tem densidade definida em (2.2). Através de (2.3) os parâmetros μ e σ^2 relacionam-se, respectivamente, com β e γ . Suponha que queremos testar a hipótese nula $H_0 : \beta_{(r)} = \beta^{(0)}$ contra a hipótese alternativa $H_1 : \beta_{(r)} \neq \beta^{(0)}$, em que $\beta = (\beta_{(r)}^\top, \beta_{(k-r)}^\top)^\top$, sendo $\beta_{(r)}$ um vetor $r \times 1$ de parâmetros de interesse e $\beta_{(k-r)}$ um vetor $(k - r) \times 1$ de parâmetros de incômodo. Dessa forma, a estatística de teste da razão de verossimilhanças generalizada é dada por

$$RV = 2\{l(\hat{\beta}, \hat{\gamma}) - l(\tilde{\beta}, \tilde{\gamma})\},$$

em que $l(\beta, \gamma)$ é o logaritmo natural da função de verossimilhança, $(\hat{\beta}^\top, \hat{\gamma}^\top)^\top$ e $(\tilde{\beta}^\top, \tilde{\gamma}^\top)^\top$ são os estimadores de máxima verossimilhança irrestrito e restrito de $(\beta^\top, \gamma^\top)^\top$, respectivamente. O estimador restrito é obtido impondo a hipótese nula.

Sejam $U_{(r)\beta}$ um vetor coluna r -dimensional contendo os r elementos iniciais da função escore de β e $K_{(r)(r)}^{\beta\beta}$ a matriz $m \times m$ formada das r linhas iniciais e das r colunas iniciais da matriz K^{-1} . A estatística escore de Rao para testar a hipótese $H_0 : \beta_{(r)} = \beta^{(0)}$ versus $H_1 : \beta_{(r)} \neq \beta^{(0)}$ é dada pela seguinte expressão

$$S_R = \tilde{U}_{(r)\beta}^\top \tilde{K}_{(r)(r)}^{\beta\beta} \tilde{U}_{(r)\beta},$$

em que $(\tilde{\cdot})$ indica as quantidades avaliadas no estimador de máxima verossimilhança restrito. No modelo de regressão simplex não linear, a estatística Wald é expressa como

$$W = (\hat{\beta}_{(r)} - \beta^{(0)})^\top (\hat{K}_{(r)(r)}^{\beta\beta})^{-1} (\hat{\beta}_{(r)} - \beta^{(0)}),$$

em que $(\hat{K}_{(r)(r)}^{\beta\beta})^{-1}$ é igual a $(K_{(r)(r)}^{\beta\beta})^{-1}$ avaliado no estimador de máxima verossimilhança irrestrito e $\hat{\beta}_{(r)}$ é o estimador de máxima verossimilhança de $\beta_{(r)}$.

Por fim, a estatística gradiente é dada por

$$S_G = \tilde{U}_{(r)\beta}^\top (\hat{\beta}_{(r)} - \beta^{(0)}).$$

Sob H_0 e considerando as condições usuais de regularidade, as estatísticas RV , S_R , W e S_G convergem em distribuição para qui-quadrado com r graus de liberdade. Dessa forma, os testes apresentados podem ser realizados com base nos valores críticos assintóticos obtidos da distribuição χ_r^2 .

4.4 TESTES DE HIPÓTESES EM REGRESSÃO SIMPLEX NÃO LINEAR SOBRE O VETOR DE PARÂMETROS γ

Suponha agora que queremos testar a hipótese nula $H_0 : \gamma_{(r)} = \gamma^{(0)}$ contra a hipótese alternativa $H_1 : \gamma_{(r)} \neq \gamma^{(0)}$, em que $\gamma = (\gamma_{(r)}^\top, \gamma_{(q-r)}^\top)^\top$, sendo $\gamma_{(r)}$ um vetor $r \times 1$ de parâmetros de interesse e $\gamma_{(q-r)}$ um vetor $(q-r) \times 1$ de parâmetros de incômodo.

Sejam $U_{(r)\gamma}$ um vetor coluna r -dimensional contendo os r elementos finais da função escore de γ e $K_{(r)(r)}^{\gamma\gamma}$ a matriz $m \times m$ formada das r linhas finais e das r colunas finais da matriz K^{-1} . As estatísticas da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W) e gradiente (S_G) são expressas, respectivamente, como

$$RV = 2\{l(\hat{\beta}, \hat{\gamma}) - l(\tilde{\beta}, \tilde{\gamma})\}, \quad S_R = \tilde{U}_{(r)\gamma}^\top \tilde{K}_{(r)(r)}^{\gamma\gamma} \tilde{U}_{(r)\gamma},$$

$$W = (\hat{\gamma}_{(r)} - \gamma^{(0)})^\top (\hat{K}_{(r)(r)}^{\gamma\gamma})^{-1} (\hat{\gamma}_{(r)} - \gamma^{(0)}) \quad \text{e} \quad S_G = \tilde{U}_{(r)\gamma}^\top (\hat{\gamma}_{(r)} - \gamma^{(0)}).$$

Assim como na seção anterior, sob H_0 e considerando as condições usuais de regularidade, as estatísticas RV , S_R , W e S_G convergem em distribuição para χ_r^2 .

4.5 TESTES DE HIPÓTESES BASEADOS EM BOOTSTRAP

Como vimos anteriormente, os testes assintóticos são válidos para grandes amostras. Quando o tamanho da amostra é pequeno ou moderado, os testes podem apresentar tamanhos bastante distorcidos. Uma alternativa para contornar esta distorção é utilizar o método bootstrap. Nesta seção, apresentamos duas versões do teste da razão de verossimilhanças generalizada baseadas neste método. A primeira se trata da correção de Bartlett bootstrap, em que o fator de correção de Bartlett é estimado via bootstrap. A segunda considera a obtenção do p -valor para o teste RV com base no bootstrap duplo rápido. A seguir, apresentaremos tais versões.

4.5.1 Correção de Bartlett Bootstrap

Com o objetivo de melhorar a estatística da razão de verossimilhanças generalizada em pequenas amostras, Bartlett (1937) modificou esta estatística por um fator de correção, conhecido na literatura como fator de correção de Bartlett, dado por $c = \mathbb{E}(RV)/r$, em que r é o número de restrições impostas pela hipótese nula. A nova versão é denominada estatística da razão de verossimilhanças modificada dada por

$$RV_B = \frac{RV}{c}.$$

Em busca de uma alternativa ao cálculo analítico do fator de correção de Bartlett, Rocke (1989) propôs utilizar o fator de correção calculado através do método bootstrap. A estatística da razão de verossimilhanças generalizada na versão Bartlett bootstrap pode ser expressa como

$$RV_{BB} = \frac{rRV}{\overline{RV^*}},$$

em que $\overline{RV^*} = \frac{1}{B} \sum_{b=1}^B RV^{*b}$ com $RV^{*b} = 2\{l(\hat{\beta}^{*b}\hat{\gamma}^{*b}; y^{*b}) - l(\tilde{\beta}^{*b}, \tilde{\gamma}^{*b}; y^{*b})\}$.

De acordo com Bayer e Cribari-Neto (2013), a correção de Bartlett bootstrap requer um número menor de reamostras comparado a correção bootstrap usual, tornando esta correção menos custosa computacionalmente. Lima e Cribari-Neto (2020) utilizaram a estatística RV_{BB} na classe de modelos de regressão beta em pequenas amostras. Os autores mostraram que em vários cenários, este teste apresenta superioridade quando comparado com o bootstrap usual. O teste RV_{BB} pode ser realizado através dos passos descritos no Algoritmo 4.

Algoritmo 4: Correção de Bartlett bootstrap

- 1: Seja $y_t \sim S^-(\mu_t, \sigma_t^2)$, $t = 1, \dots, n$, com $\mu_t = g^{-1}(f_1(x_t^\top; \beta))$ e $\sigma_t^2 = h^{-1}(f_2(z_t^\top; \gamma))$ uma amostra e Φ um teste de hipótese de interesse sobre os vetores β e γ ;
- 2: Calcule o valor da estatística de teste RV para a realização de Φ com base nesta amostra;
- 3: Gere B amostras bootstrap de tamanho n , a saber: $(y_1^*, \dots, y_n^*)$, $y_t^* \sim S^-(\tilde{\mu}_t, \tilde{\sigma}_t^2)$ com $\tilde{\mu}_t = g^{-1}(f_1(x_t^\top; \tilde{\beta}))$ e $\tilde{\sigma}_t^2 = h^{-1}(f_2(z_t^\top; \tilde{\gamma}))$, em que $\tilde{\beta}$ e $\tilde{\gamma}$ são as estimativas restritas de MV de β e γ obtidas pela imposição da hipótese nula de Φ , respectivamente;
- 4: Para cada amostra bootstrap calcule a estatística RV^{*b} correspondente, com $b = 1, \dots, B$, tal que

$$RV^{*b} = 2\{l(\hat{\beta}^{*b}\hat{\gamma}^{*b}; y^{*b}) - l(\tilde{\beta}^{*b}, \tilde{\gamma}^{*b}; y^{*b})\};$$

- 5: Calcule a média $\overline{RV^*}$ e em seguida a estatística corrigida RV_{BB} ;
 - 6: Rejeite a hipótese nula se $RV_{BB} > \chi_{1-\alpha, r}^2$.
-

4.5.2 Teste Bootstrap Duplo Rápido

O teste bootstrap duplo foi proposto por Efron (1983), mas devido ao alto custo computacional Davidson e MacKinnon (2000) propuseram o teste bootstrap duplo rápido (BDR). Pesquisadores como Davidson e MacKinnon (2002), MacKinnon (2006) e Davidson e MacKinnon (2007) mostram que em pequenas amostras o teste BDR apresenta bons resultados quando comparado com o bootstrap tradicional e o bootstrap duplo. Recentemente, Lima e Cribari-Neto (2020) propuseram a versão do teste bootstrap duplo rápido para a realização de testes de hipóteses na classe de modelos de regressão beta em pequenas amostras. Os autores recomendam a utilização deste teste, visto que o mesmo apresenta superioridade quando comparado com as demais versões bootstrap utilizadas, além de apresentar baixo custo computacional.

Sejam $y = (y_1, \dots, y_n)^\top$ uma amostra aleatória de $y \sim F(y, \theta)$, Φ um teste de hipótese de interesse sobre θ , τ uma estatística de teste e $\hat{\tau}$ o valor que a mesma assume para a amostra y . Com base em y , impondo H_0 de Φ e a partir de $F_{\hat{\theta}}$ geramos B pseudo-amostras y^{*b} , obtendo as estatísticas $\tau^{*b} = \tau(y^{*b})$, $b = 1, \dots, B$. Posteriormente, para cada y^{*b} impondo H_0 e a partir de $F_{\hat{\theta}^{*b}}$ geramos uma amostra bootstrap de segundo nível y_b^{**} , cuja estatística é $\tau_b^{**} = \tau(y_b^{**})$. O p -valor BDR é dado por

$$p_{DR}^{**}(\hat{\tau}) = \frac{1}{B} \sum_{b=1}^B I(\tau^{*b} > \hat{Q}_B^{**}(1 - p^*(\hat{\tau}))), \quad (4.2)$$

em que $\hat{Q}_B^{**}(1 - p^*(\hat{\tau}))$ é o quantil $1 - p^*(\hat{\tau})$ de τ_b^{**} e

$$p^*(\hat{\tau}) = \frac{1}{B} \sum_{b=1}^B I(\tau^{*b} > \hat{\tau}),$$

em que $I(\cdot)$ é a função indicadora.

Segundo MacKinnon (2006), quando $B \rightarrow \infty$, a Equação (4.2) pode ser expressa por

$$p_{DR}^{**}(\hat{\tau}) = Pr_*(\tau^{*b} > Q^{**}(1 - p^*)),$$

em que $Q^{**}(x) = \lim_{B \rightarrow \infty} \hat{Q}_B^{**}(x)$ e o subscrito $*$ indica que a probabilidade foi tomada com respeito à distribuição dos τ^{*b} .

Os passos para a construção do teste bootstrap duplo rápido estão descritos no Algoritmo 5.

Algoritmo 5: Bootstrap duplo rápido

- 1: Seja $y_t \sim S^-(\mu_t, \sigma_t^2)$, $t = 1, \dots, n$, com $\mu_t = g^{-1}(f_1(x_t^\top; \beta))$ e $\sigma_t^2 = h^{-1}(f_2(z_t^\top; \gamma))$ uma amostra e Φ um teste de hipótese de interesse sobre os vetores β e γ ;
- 2: Calcule o valor da estatística de teste RV para a realização de Φ com base nesta amostra;
- 3: Gere uma amostra bootstrap y^* , sob H_0 , com $y_t^* \sim S^-(\tilde{\mu}_t, \tilde{\sigma}_t^2)$, em que $\tilde{\mu}_t = g^{-1}(f_1(x_t^\top; \tilde{\beta}))$, $\tilde{\sigma}_t^2 = h^{-1}(f_2(z_t^\top; \tilde{\gamma}))$ e $\tilde{\beta}$ e $\tilde{\gamma}$ são as estimativas restritas de MV de β e γ obtidas pela imposição da hipótese nula de Φ , respectivamente;
- 4: Calcule a estatística RV^* associada a amostra bootstrap

$$RV^* = 2\{l(\hat{\beta}^*, \hat{\gamma}^*; y^*) - l(\tilde{\beta}^*, \tilde{\gamma}^*; y^*)\};$$

- 5: Gere uma amostra bootstrap de segundo nível y^{**} , sob H_0 , com $y_t^{**} \sim S^-(\tilde{\mu}_t^*, \tilde{\sigma}_t^{*2})$, em que $\tilde{\mu}_t^* = g^{-1}(f_1(x_t^\top; \tilde{\beta}^*))$, $\tilde{\sigma}_t^{*2} = h^{-1}(f_2(z_t^\top; \tilde{\gamma}^*))$ e $\tilde{\beta}^*$ e $\tilde{\gamma}^*$ são as estimativas restritas obtidas pela imposição da hipótese nula de β e γ , respectivamente, utilizando y^* como variável resposta;
- 6: Obtenha a estatística RV^{**} associada à amostra bootstrap de segundo nível, como no passo 4, por exemplo;
- 7: Faça os passos 3 a 6 um número grande (B) de vezes;
- 8: Calcule o p -valor de primeiro nível, expresso como

$$p^*(RV) = \frac{1}{B} \sum_{b=1}^B I(RV^{*b} > RV);$$

- 9: Encontre o quantil $1 - p^*$ de RV_b^{**} ($b = 1, \dots, B$), $\hat{Q}_B^{**}(1 - p^*(RV))$;
- 10: Obtenha o p -valor bootstrap duplo rápido, dado por

$$p_{DR}^{**}(RV) = \frac{1}{B} \sum_{b=1}^B I(RV^{*b} > \hat{Q}_B^{**}(1 - p^*(RV)));$$

- 11: Rejeita-se a hipótese nula H_0 se p_{DR}^{**} é menor que o nível nominal adotado.
-

4.6 RESULTADOS NUMÉRICOS

Nesta seção apresentamos os resultados de simulações de Monte Carlo realizadas para avaliar os desempenhos dos testes da razão de verossimilhanças generalizada, escore, Wald, gradiente, razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap e da razão de verossimilhanças generalizada bootstrap duplo rápido em pequenas amostras. Nas simulações foram consideradas $R = 5000$ réplicas de Monte Carlo e $B = 500$ réplicas bootstrap. As estimativas dos parâmetros dos modelos que apresentaremos foram obtidas por maximização da função de log-verossimilhança usando o método de otimização não linear Escore de Fisher. Os resultados numéricos foram obtidos por meio da linguagem de programação

Ox (DOORNIK, 2001; DOORNIK, 2006; DOORNIK, 2009).

4.6.1 Cenário 1

No Cenário 1 vamos considerar o seguinte modelo de regressão simplex não linear:

$$g(\mu_t) = \beta_1 + x_{t2}^{\beta_2} + \beta_3 x_{t3} + \beta_4 x_{t4} \quad \text{e} \quad h(\sigma_t^2) = \gamma_1 + z_{t2}^{\gamma_2} + \gamma_3 z_{t3}, \quad t = 1, \dots, n,$$

em que $g(\cdot)$ e $h(\cdot)$ são as funções de ligação logito e logarítmica, respectivamente.

As realizações das covariáveis foram geradas através da distribuição uniforme como segue: $x_{t2} \sim \mathcal{U}(0.5, 1.5)$, $x_{t3} \sim \mathcal{U}(0, 1)$, $x_{t4} \sim \mathcal{U}(-0.5, 0.5)$, $z_{t2} \sim \mathcal{U}(0.5, 1.5)$ e $z_{t3} \sim \mathcal{U}(-0.5, 0.5)$ e mantidas fixas para cada réplica de Monte Carlo. Deseja-se testar $H_0 : \beta_3 = \beta_4 = 0$ ($r = 2$), em que r é o número de restrição da hipótese nula. Foram considerados três diferentes cenários para a resposta média: $\mu_t \in (0.04, 0.13)$ com $\beta = (-3.5, 1.2, 0, 0)^\top$; $\mu_t \in (0.19, 0.82)$ com $\beta = (-1.9, -1.8, 0, 0)^\top$ e $\mu_t \in (0.92, 0.98)$ com $\beta = (2.0, 1.4, 0, 0)^\top$. Além disso, para medir a intensidade de dispersão não constante, consideramos $\lambda = \sigma_{\max}^2 / \sigma_{\min}^2$ e reportamos os resultados para $\lambda \approx 9$ com $\gamma = (-1.2, 2.2, 1.3)^\top$; $\lambda \approx 53$ com $\gamma = (-1.2, -1.9, 1.4)^\top$ e $\lambda \approx 101$ com $\gamma = (-1.2, -2.2, 1.3)^\top$. Os tamanhos amostrais considerados foram $n = 15, 30$ e 45 . Para os dois últimos casos foram geradas inicialmente $n = 15$ observações das covariáveis, sendo estas replicadas duas e três vezes, respectivamente, para obter os tamanhos amostrais $n = 30$ e $n = 45$. Isso foi feito para garantir que a intensidade de dispersão não constante fosse a mesma para todos os tamanhos amostrais.

As Tabelas 14, 15 e 16 consideram, respectivamente, $\lambda \approx 9$, $\lambda \approx 53$ e $\lambda \approx 101$. Nestas tabelas são apresentadas as taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), score (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras para testar a hipótese $H_0 : \beta_3 = \beta_4 = 0$.

No geral, conforme aumenta-se o tamanho da amostra a taxa de rejeição nula fica mais próxima do nível nominal adotado. O teste da razão de verossimilhanças generalizada apresenta comportamento liberal. Por exemplo, para $\lambda \approx 53$, $\mu_t \in (0.04, 0.13)$, $\alpha = 10\%$ e $n = 45$, a taxa de rejeição nula do teste foi 15.58%.

O teste de Wald foi o que apresentou pior desempenho com distorções de tamanho muito elevadas, sendo acentuadamente liberal. Por exemplo, para o caso em que $\lambda \approx 9$, $\mu_t \in$

$(0.04, 0.13)$, $\alpha = 10\%$ e $n = 15$, o teste rejeita a hipótese nula em cerca de 60% das vezes. Quando aumentamos o tamanho da amostra para $n = 45$, o teste rejeita H_0 cerca de 20% das vezes.

Percebe-se, no geral, que à medida que aumentamos a intensidade de dispersão não constante os desempenhos dos testes RV e W tendem a piorar. Por exemplo, para o teste RV , $\lambda \approx 9$, $\mu_t \in (0.04, 0.13)$, $\alpha = 10\%$ e $n = 15$, a taxa de rejeição nula foi 38.56%, quando λ passa a ser $\lambda \approx 101$, a taxa foi 47.74%.

Em todos os cenários considerados, o teste escore apresenta desempenho superior aos testes RV e Wald. Por exemplo, na Tabela 16 em que $\lambda \approx 101$, para $\mu_t \in (0.19, 0.82)$, $\alpha = 5\%$ e $n = 45$, a taxa de rejeição nula do teste escore foi 5.68%, enquanto que as dos testes RV e Wald foram 9.42% e 13.12%, respectivamente. Além disso, nota-se que o teste escore apresenta desempenho superior ao teste gradiente em todos os cenários em que $\alpha = 10\%$ e 5%. Porém, quando $\alpha = 1\%$ o teste gradiente se saiu melhor. Por exemplo, para $\lambda \approx 53$, $\mu_t \in (0.19, 0.82)$, $\alpha = 1\%$ e $n = 30$, a taxa de rejeição nula do teste gradiente foi 0.94%, enquanto que a do teste escore 0.70%. No geral, o teste escore apresenta bom desempenho comparado com os testes RV , Wald e gradiente. Os tamanhos dos testes escore e gradiente são pouco afetados pelo aumento da intensidade de dispersão não constante. Notamos que quando o nível nominal é igual a 1%, o teste escore é tipicamente conservador, ou seja, apresenta tamanho menor que o nível nominal de referência. Para $\lambda \approx 53$, $\mu_t \in (0.92, 0.98)$, $\alpha = 1\%$ e $n = 45$, os testes escore, Wald e gradiente apresentaram comportamentos acentuadamente liberais, o que não ocorre nos demais cenários em que $\alpha = 1\%$ e $n = 45$.

Nota-se que o teste gradiente é bastante liberal para o caso em que $n = 15$ e níveis de significância $\alpha = 10\%$ e $\alpha = 5\%$, pois o teste apresenta grande distorção de tamanho com taxa de rejeição nula longe do nível nominal. Por exemplo, para $\lambda \approx 101$, $\mu_t \in (0.04, 0.13)$, $\alpha = 10\%$ e $n = 15$, o teste apresenta tamanho igual a 21.08%, enquanto que o do teste escore 13.26%.

Ainda em relação ao controle da probabilidade de erro tipo I, os testes baseados em reamostragem bootstrap apresentam melhor desempenho que os demais testes. Fazendo um comparativo entre o teste RV e os testes baseados nessa estatística, percebemos que o método bootstrap reduz consideravelmente a distorção de tamanho. Por exemplo, na Tabela 14, quando $\mu_t \in (0.19, 0.82)$, $\alpha = 10\%$ e $n = 15$, a taxa de rejeição nula do teste RV foi 35.88%, enquanto que as versões RV_{BB} e RV_{BDR} apresentaram taxas de rejeição nula iguais a 9.72% e 10.16%, respectivamente. Nota-se que o teste RV_{BB} tem desempenho equivalente ao teste

RV_{BDR} . Esses testes apresentam melhor desempenho no controle da probabilidade de erro tipo I.

Portanto, o teste RV apresenta desempenho inferior aos testes escore e gradiente, e apesar de ter um desempenho pobre em relação ao controle da probabilidade de erro tipo I, apresentou taxas de rejeição nula menos distorcidas que o teste de Wald. Este por sua vez, mostrou comportamento acentuadamente liberal em pequenas amostras. O teste escore apresentou bom desempenho. O teste gradiente mostrou desempenho pobre para $\alpha = 10\%$ e 5% , mas para $\alpha = 1\%$ a taxa de rejeição nula ficou mais próxima do nível desejado. Por fim, os testes baseados em reamostragem bootstrap apresentaram desempenho superior a todos os demais testes.

Tabela 14 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_3 = \beta_4 = 0$ ($r = 2$). $\lambda \approx 9$.

$\mu_t \in (0.04, 0.13)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	38.56	18.22	15.18	29.04	11.04	8.14	14.86	3.70	2.12
S_R	12.72	11.48	10.96	5.68	5.66	4.86	0.50	0.86	0.70
W	60.26	26.48	20.30	53.18	18.64	12.58	42.82	9.44	4.50
S_G	20.46	15.14	13.36	11.10	7.46	6.08	1.74	1.34	1.16
RV_{BB}	9.08	10.16	10.10	4.62	5.32	4.86	0.68	1.10	1.00
RV_{BDR}	9.32	10.36	10.26	4.96	5.40	5.04	1.04	1.28	1.06
$\mu_t \in (0.19, 0.82)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	35.88	17.98	14.60	27.34	10.22	8.38	14.02	2.92	2.32
S_R	13.30	10.58	10.92	5.42	4.90	5.50	0.52	0.74	0.84
W	54.42	25.48	18.86	47.14	18.62	12.56	37.18	8.60	4.98
S_G	21.74	14.50	12.82	11.24	7.20	6.80	1.60	1.34	1.30
RV_{BB}	9.72	9.86	10.30	4.32	5.04	5.52	0.80	1.00	1.18
RV_{BDR}	10.16	10.20	10.60	4.66	4.94	5.64	0.96	0.96	1.18
$\mu_t \in (0.92, 0.98)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	35.32	17.48	15.04	26.00	10.60	8.46	12.80	2.78	2.08
S_R	14.04	11.80	11.40	6.22	5.22	5.44	0.30	0.72	0.74
W	53.90	24.24	18.70	46.84	16.76	12.04	36.64	8.16	4.30
S_G	20.52	14.22	13.24	11.14	7.30	6.68	1.44	1.18	1.12
RV_{BB}	9.58	10.08	10.78	4.48	4.94	5.44	0.84	1.00	0.94
RV_{BDR}	9.58	10.26	10.74	4.80	4.98	5.48	1.16	1.20	1.04

Fonte: próprio autor.

Tabela 15 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_3 = \beta_4 = 0$ ($r = 2$). $\lambda \approx 53$.

$\mu_t \in (0.04, 0.13)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	45.10	18.82	15.58	35.20	11.54	9.14	19.14	4.02	2.56
S_R	12.16	10.68	10.84	4.84	4.96	5.48	0.28	0.58	0.80
W	70.46	28.06	20.72	64.98	20.32	13.60	55.06	10.78	5.78
S_G	20.20	13.14	12.30	9.50	6.18	6.28	2.00	0.90	0.92
RV_{BB}	8.96	9.44	9.98	3.72	4.90	5.52	0.56	1.08	0.94
RV_{BDR}	9.30	9.56	10.16	4.42	5.12	5.26	1.02	1.28	1.08

$\mu_t \in (0.19, 0.82)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	40.04	19.14	14.46	29.90	12.14	7.90	15.24	3.50	2.12
S_R	13.32	11.44	9.88	5.94	5.16	4.92	0.42	0.70	0.62
W	62.18	27.24	19.22	55.74	19.46	11.82	44.74	10.08	4.56
S_G	21.12	14.78	11.62	10.44	7.00	5.72	1.70	0.94	0.78
RV_{BB}	10.74	10.72	9.44	5.04	5.38	4.62	0.92	1.02	0.84
RV_{BDR}	10.00	10.64	9.36	5.02	5.42	4.82	0.82	1.00	0.98

$\mu_t \in (0.92, 0.98)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	41.96	19.10	15.16	32.80	11.30	8.30	16.82	3.48	2.50
S_R	13.34	10.84	11.04	5.10	4.52	5.36	0.40	0.72	11.04
W	64.78	28.58	19.86	58.68	20.44	12.74	48.78	10.34	19.86
S_G	19.84	13.36	12.12	9.02	6.22	5.84	1.58	0.80	12.12
RV_{BB}	10.20	9.32	9.80	4.10	4.50	5.06	0.58	0.78	9.80
RV_{BDR}	10.72	9.54	9.58	5.16	4.54	4.92	1.00	0.94	9.58

Fonte: próprio autor.

4.6.2 Cenário 2

Dando continuidade à avaliação dos testes, foi considerado o modelo com parâmetro de dispersão constante, função de ligação probito e uma nova expressão do preditor não linear para o submodelo da média, a saber:

$$g(\mu_t) = \beta_1 + \exp\{\beta_2 x_{t2}\} + \beta_3 x_{t3}, \quad t = 1, \dots, n.$$

Antes de definirmos a composição dos parâmetros das simulações vamos ressaltar que a nova definição da função não linear para o submodelo da média, nos permite a realização do teste em que a hipótese nula é dada por $H_0 : \beta_2 = 1.0$, o que implica linearidade do modelo. Dado o interesse em avaliar os desempenhos dos testes quanto ao teste de linearidade do modelo, colocamos esse cenário para o preditor não linear do submodelo da média.

Tabela 16 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_3 = \beta_4 = 0$ ($r = 2$). $\lambda \approx 101$.

$\mu_t \in (0.04, 0.13)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	47.74	19.32	15.08	37.94	11.84	8.82	20.40	3.74	2.52
S_R	13.26	10.62	10.62	5.32	4.70	4.88	0.36	0.54	0.66
W	71.12	29.26	20.16	65.94	20.90	12.84	56.76	10.84	5.60
S_G	21.08	13.40	12.08	10.20	6.28	5.68	1.70	0.78	0.74
RV_{BB}	9.72	9.92	9.90	4.20	5.04	5.14	0.52	0.86	1.04
RV_{BDR}	9.94	10.00	9.92	5.38	4.90	5.06	1.14	1.04	1.00
$\mu_t \in (0.19, 0.82)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	40.78	18.56	16.24	31.52	11.20	9.42	16.52	3.34	2.42
S_R	13.82	11.08	11.96	5.48	4.78	5.68	0.32	0.70	0.78
W	61.30	26.14	20.72	54.80	19.28	13.12	44.80	9.16	5.46
S_G	22.02	13.54	13.54	10.54	6.38	6.78	1.78	0.92	1.02
RV_{BB}	11.68	9.82	11.30	6.16	4.84	5.82	1.48	1.00	1.00
RV_{BDR}	10.22	9.80	11.06	5.86	4.98	6.00	1.22	1.02	1.18
$\mu_t \in (0.92, 0.98)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	43.12	20.04	15.44	32.52	12.06	8.64	16.60	3.86	2.16
S_R	13.20	11.90	11.12	5.60	5.58	5.02	0.32	0.58	0.68
W	66.08	28.72	20.86	59.32	20.96	13.04	49.38	10.82	4.96
S_G	20.10	13.94	12.70	9.76	6.72	5.82	1.44	0.90	0.78
RV_{BB}	9.58	10.52	10.24	4.44	5.02	4.82	0.66	1.00	0.98
RV_{BDR}	10.24	10.80	10.26	4.90	5.18	5.20	1.16	1.04	0.92

Fonte: próprio autor.

As realizações das covariáveis foram geradas através da distribuição uniforme como segue: $x_{t2} \sim \mathcal{U}(-0.5, 0.5)$ e $x_{t3} \sim \mathcal{U}(0.5, 1.5)$ e mantidas fixas para cada réplica de Monte Carlo. Deseja-se testar $H_0 : \beta_2 = 1$ ($r = 1$). Foram considerados três diferentes cenários para a resposta média: $\mu_t \in (0.01, 0.12)$ com $\beta = (-3.3, 1.0, 0.4)^\top$; $\mu_t \in (0.14, 0.80)$ com $\beta = (-2.9, 1.0, 1.5)^\top$ e $\mu_t \in (0.84, 0.99)$ com $\beta = (1.5, 1.0, -1.3)^\top$. Consideramos quatro valores distintos para o parâmetro de dispersão, a saber: $\sigma^2 = 0.5, 1.5, 3.0$ e 6.0 .

As Tabelas 17, 18, 19 e 20 consideram, respectivamente, $\sigma_2 = 0.5, 1.5, 3.0$ e 6.0 . Nestas tabelas são apresentadas as taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada, escore, Wald, gradiente, razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap e da razão de verossimilhanças generalizada bootstrap duplo rápido em pequenas amostras para testar a hipótese $H_0 : \beta_2 = 1$.

No geral, conforme aumenta-se o tamanho da amostra a taxa de rejeição nula aproxima-se

do nível nominal adotado. O teste da razão de verossimilhanças generalizada não apresenta bom desempenho. Por exemplo, na Tabela 18, quando $\mu_t \in (0.84, 0.99)$, $\alpha = 10\%$ e $n = 15$, a taxa de rejeição nula do teste foi 15.68%. Para $n = 30$ e 45, as taxas foram 11.82% e 11.44%, respectivamente, ou seja, se aproximam do nível desejado mas ainda ficam distantes de 10%.

O teste de Wald foi o que apresentou pior desempenho com distorções de tamanho muito elevadas. Por exemplo, na Tabela 20, quando $\mu_t \in (0.14, 0.80)$, $\alpha = 5\%$ e $n = 15$, o teste rejeita a hipótese nula em cerca de 20% das vezes, quando $n = 45$ o teste rejeita em cerca de 14% das vezes.

Dos testes assintóticos, o escore foi o que apresentou melhor desempenho, ou seja, foi o que apresentou taxas de rejeição nula mais próximas do nível nominal de referência. Por exemplo, na Tabela 17, quando $\mu_t \in (0.01, 0.12)$, $\alpha = 1\%$ e $n = 45$, a taxa de rejeição nula do teste escore foi 1.10%, enquanto que as dos testes RV , W e S_G foram 1.38%, 1.82% e 1.14%, respectivamente.

Ainda em relação aos testes assintóticos, o teste gradiente foi o segundo com melhor desempenho. O mesmo apresentou taxas de rejeição nula mais próximas do nível desejado. Em alguns casos, ele se saiu melhor que o teste escore. Por exemplo, na Tabela 19, quando $\mu_t \in (0.14, 0.80)$, $\alpha = 1\%$ e $n = 15, 30$ e 45, as taxas do teste gradiente foram 1.02%, 0.92% e 1.02%, enquanto que as do teste escore foram 0.84%, 0.84% e 0.90%, respectivamente.

Os testes baseados em reamostragem bootstrap foram os que apresentaram melhor desempenho. Em alguns casos o teste RV_{BB} se saiu melhor. Em outros, o teste RV_{BDR} . No geral, os dois testes se saíram muito bem. Por exemplo, na Tabela 19, quando $\mu_t \in (0.84, 0.99)$, $\alpha = 5\%$ e $n = 30$, o teste RV_{BB} apresenta probabilidade de erro tipo I de 5.02%, enquanto que a do teste RV_{BDR} de 5.06%.

Percebe-se que os tamanhos dos testes não foram muito afetados com o aumento dos valores considerados para o parâmetro de dispersão σ^2 . Os testes apresentaram tamanhos bastante similares. Dentre os testes assintóticos tradicionais, o teste escore foi o que apresentou melhor desempenho com taxas de rejeição nula mais próximas do nível nominal desejado. Adicionalmente, os testes baseados em reamostragem bootstrap mostraram ser eficientes no controle da probabilidade de erro tipo I.

Tabela 17 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_2 = 1$ ($r = 1$). $\sigma^2 = 0.5$.

$\mu_t \in (0.01, 0.12)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	15.04	12.22	11.20	8.54	6.92	6.12	2.34	1.72	1.38
S_R	12.98	11.34	10.68	6.30	5.92	5.50	0.92	1.10	1.10
W	17.16	13.28	12.00	10.56	7.80	6.58	3.88	2.38	1.82
S_G	13.02	11.38	10.72	6.42	6.00	5.60	0.88	1.22	1.14
RV_{BB}	10.06	10.22	10.22	5.04	5.24	5.14	1.02	1.24	1.06
RV_{BDR}	10.26	10.44	10.40	5.32	5.22	4.96	0.98	1.26	0.90
$\mu_t \in (0.14, 0.80)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	14.82	11.64	11.26	8.62	6.64	5.92	2.40	1.58	1.18
S_R	12.48	10.82	10.56	6.22	5.68	5.36	0.86	1.04	0.92
W	17.00	12.66	11.80	10.60	7.34	6.38	4.20	2.10	1.58
S_G	13.02	10.74	10.70	6.34	5.66	5.40	0.82	0.94	0.92
RV_{BB}	9.98	9.70	10.00	4.78	5.08	4.96	0.92	0.98	0.94
RV_{BDR}	9.90	9.90	10.22	5.04	4.88	5.14	1.20	1.02	0.98
$\mu_t \in (0.84, 0.99)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	15.96	12.34	12.16	9.28	6.60	6.54	2.90	1.48	1.46
S_R	14.00	11.58	11.64	6.98	5.36	5.84	1.06	0.92	1.12
W	17.56	13.24	12.64	11.54	7.48	7.10	4.50	2.02	1.86
S_G	14.00	11.64	11.68	7.00	5.44	5.92	1.10	0.86	1.12
RV_{BB}	10.92	10.38	10.92	5.58	5.06	5.54	1.12	0.84	1.18
RV_{BDR}	10.76	10.14	10.62	5.74	4.96	5.76	1.16	0.96	1.38

Fonte: próprio autor.

4.6.3 Cenário 3

Neste cenário o preditor do submodelo da média é o mesmo do Cenário 2. No entanto, vamos considerar a modelagem da dispersão, sendo que o preditor será linear como descrevemos a seguir:

$$g(\mu_t) = \beta_1 + \exp\{\beta_2 x_{t2}\} + \beta_3 x_{t3} \quad \text{e} \quad h(\sigma_t^2) = \gamma_1 + \gamma_2 z_{t2} + \gamma_3 z_{t3}, \quad t = 1, \dots, n,$$

em que $g(\cdot)$ e $h(\cdot)$ são as funções de ligação probito e logarítmica, respectivamente.

As realizações das covariáveis foram geradas através da distribuição uniforme como segue: $x_{t2} \sim \mathcal{U}(-0.5, 0.5)$, $x_{t3} \sim \mathcal{U}(0.5, 1.5)$, $z_{t2} \sim \mathcal{U}(0.5, 1.5)$ e $z_{t3} \sim \mathcal{U}(0, 1)$ e mantidas fixas para cada réplica de Monte Carlo. A hipótese nula considerada foi $H_0 : \gamma_2 = \gamma_3 = 0$ ($r = 2$). Os testes foram realizados com base nos seguintes cenários: $\mu_t \in (0.01, 0.13)$ com $\beta =$

Tabela 18 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_2 = 1$ ($r = 1$). $\sigma^2 = 1.5$.

$\mu_t \in (0.01, 0.12)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	15.50	12.00	11.36	9.02	6.08	5.64	2.76	1.30	1.26
S_R	12.64	10.86	10.60	6.50	5.10	4.86	1.26	0.80	0.84
W	17.42	13.06	12.20	11.30	7.16	6.54	4.70	1.86	1.58
S_G	13.28	11.06	10.76	6.88	5.08	5.10	1.14	0.68	0.78
RV_{BB}	10.30	9.98	9.94	5.50	4.58	4.68	1.18	0.78	0.92
RV_{BDR}	10.24	9.58	9.96	5.52	4.82	4.88	1.30	0.94	0.92
$\mu_t \in (0.14, 0.80)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	15.18	12.46	11.30	8.94	6.82	5.90	3.00	1.58	1.56
S_R	12.66	11.40	10.60	6.30	5.60	5.14	1.22	0.88	1.18
W	17.82	13.84	12.26	12.02	8.10	6.76	4.96	2.44	2.18
S_G	13.20	11.66	10.76	7.06	5.92	5.30	1.18	1.08	1.20
RV_{BB}	10.28	10.50	10.02	5.28	5.26	5.06	1.34	0.94	1.20
RV_{BDR}	10.14	10.40	10.00	5.34	5.26	5.24	1.38	1.08	1.18
$\mu_t \in (0.84, 0.99)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	15.68	11.82	11.44	9.44	6.42	5.78	2.50	1.58	1.28
S_R	13.58	10.74	10.50	6.76	5.50	4.98	1.00	0.96	0.84
W	17.78	12.70	12.08	11.44	7.28	6.48	4.20	2.22	1.64
S_G	13.76	10.86	10.86	7.06	5.66	5.00	1.02	1.10	0.94
RV_{BB}	10.84	9.32	9.98	5.62	5.14	4.74	1.10	1.12	0.92
RV_{BDR}	10.76	9.44	9.96	5.40	5.14	5.06	1.08	1.06	0.88

Fonte: próprio autor.

$(-1.9, 0.8, -1.1)^\top$; $\mu_t \in (0.15, 0.83)$ com $\beta = (-2.0, 1.7, 0.6)^\top$ e $\mu_t \in (0.87, 0.98)$ com $\beta = (1.4, 0.8, -1.1)^\top$. Além disso, foram considerados quatro diferentes cenários para o parâmetro de dispersão, a saber: $\sigma^2 = 0.5$ com $\gamma = (-0.7, 0, 0)^\top$; $\sigma^2 = 1.5$ com $\gamma = (0.4, 0, 0)^\top$, $\sigma^2 = 3.0$ com $\gamma = (1.1, 0, 0)^\top$ e $\sigma^2 = 6.0$ com $\gamma = (1.8, 0, 0)^\top$.

As Tabelas 21, 22, 23 e 24 apresentam as taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada, escore, Wald, gradiente, razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap e da razão de verossimilhanças generalizada bootstrap duplo rápido em pequenas amostras para testar a hipótese $H_0 : \gamma_2 = \gamma_3 = 0$ considerando, respectivamente, $\sigma^2 = 0.5$, $\sigma^2 = 1.5$, $\sigma^2 = 3.0$ e $\sigma^2 = 6.0$.

Note que o teste RV apresenta comportamento bastante liberal. Por exemplo, na Tabela 21, quando $\mu_t \in (0.01, 0.13)$, $\alpha = 10\%$ e $n = 15$, a taxa de rejeição nula do teste foi 31.86%. Aumentando o tamanho da amostra para $n = 45$ a taxa passa a ser 13.82%. O melhor

Tabela 19 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_2 = 1$ ($r = 1$). $\sigma^2 = 3.0$.

$\mu_t \in (0.01, 0.12)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	14.84	11.62	11.62	8.74	5.94	5.86	2.70	1.42	1.26
S_R	12.26	10.54	10.98	6.12	5.12	4.98	1.10	1.16	0.98
W	16.76	13.08	12.06	11.06	7.14	6.94	4.76	2.14	1.58
S_G	12.64	10.78	11.10	6.36	5.16	5.16	1.06	1.04	0.96
RV_{BB}	10.02	9.48	10.28	5.10	4.52	4.90	1.12	1.04	0.92
RV_{BDR}	9.86	9.52	10.12	5.20	4.46	4.78	1.16	0.98	1.12
$\mu_t \in (0.14, 0.80)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	15.32	12.46	11.62	9.22	6.92	6.14	2.50	1.48	1.34
S_R	12.38	10.74	10.98	6.22	5.56	5.44	0.84	0.84	0.90
W	18.38	14.06	12.60	12.78	8.52	7.32	5.58	2.56	2.02
S_G	13.56	11.26	11.04	7.10	5.76	5.64	1.02	0.92	1.02
RV_{BB}	10.66	9.96	10.56	5.12	5.12	5.12	0.92	0.92	1.10
RV_{BDR}	10.58	10.12	10.18	5.20	5.08	5.00	0.96	0.96	1.24
$\mu_t \in (0.84, 0.99)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	15.38	12.64	10.98	9.12	6.74	5.72	2.66	1.54	1.32
S_R	12.56	10.88	10.18	6.18	5.58	4.82	1.08	1.00	1.02
W	17.72	13.34	11.64	11.48	7.80	6.40	4.32	2.24	1.76
S_G	13.06	11.48	10.38	6.66	5.64	5.14	0.86	0.96	1.00
RV_{BB}	10.46	10.02	9.58	5.22	5.02	4.82	1.10	1.12	1.02
RV_{BDR}	10.42	10.18	9.72	5.44	5.06	5.12	1.20	1.06	1.02

Fonte: próprio autor.

desempenho deste teste ocorre quando $\alpha = 1\%$ e $n = 45$. Por exemplo, para o caso em que $\sigma^2 = 0.5$, $\mu_t \in (0.01, 0.13)$, $\alpha = 1\%$ e $n = 45$, o teste apresenta taxa de rejeição de 1.84%.

No geral, o teste escore apresenta comportamento conservador, ou seja, apresenta tamanho menor que o nível nominal desejado. No entanto, este foi o que apresentou melhor desempenho entre os testes assintóticos. Por exemplo, na Tabela 22, quando $\mu_t \in (0.87, 0.98)$, $\alpha = 5\%$ e $n = 45$, a taxa de rejeição nula do teste foi 4.86%, enquanto que as dos testes RV , Wald e gradiente foram 7.40%, 12.18% e 6.96%, respectivamente.

O teste de Wald foi o que apresentou pior desempenho no controle da probabilidade estimada de erro tipo I, sendo acentuadamente liberal. Por exemplo, para o caso em que $\sigma^2 = 3.0$, $\mu_t \in (0.01, 0.13)$, $\alpha = 10\%$ e $n = 15$, o teste rejeita a hipótese nula em cerca de 60% das vezes. E mesmo aumentando o tamanho da amostra para $n = 45$ o teste apresenta grande distorção de tamanho com taxa de rejeição nula de 21.04%.

Tabela 20 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \beta_2 = 1$ ($r = 1$). $\sigma^2 = 6.0$.

$\mu_t \in (0.01, 0.12)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	15.80	12.48	11.28	9.24	6.58	5.74	2.86	1.70	1.46
S_R	11.92	10.68	10.12	6.22	5.26	5.02	1.24	0.88	1.24
W	19.10	13.90	12.72	12.76	8.26	6.90	4.98	2.70	2.08
S_G	13.16	11.42	10.44	6.42	5.52	5.14	1.14	0.94	1.12
RV_{BB}	10.42	10.16	9.76	5.14	4.96	4.98	1.36	1.24	1.14
RV_{BDR}	10.32	10.08	9.70	5.20	4.86	5.00	1.52	1.14	1.32
$\mu_t \in (0.14, 0.80)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	15.36	13.30	12.10	8.80	7.14	6.28	2.32	1.66	1.48
S_R	11.00	11.22	10.52	5.38	5.64	5.36	0.82	0.90	0.92
W	19.60	15.32	13.98	13.50	9.64	7.88	6.28	3.14	2.40
S_G	12.94	12.48	11.44	6.56	6.14	5.66	1.24	1.30	1.14
RV_{BB}	10.18	11.06	10.54	4.84	5.44	5.42	1.10	1.02	1.24
RV_{BDR}	10.30	10.84	10.26	4.96	5.44	5.20	1.06	1.08	1.24
$\mu_t \in (0.84, 0.99)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	14.00	11.72	11.24	7.62	6.46	5.80	2.42	1.46	1.24
S_R	10.22	10.18	10.38	4.90	5.06	4.68	0.80	0.76	0.80
W	16.34	12.88	12.04	10.16	7.64	6.62	4.06	2.02	1.86
S_G	11.34	10.94	10.44	5.32	5.50	5.10	0.74	0.90	0.96
RV_{BB}	8.72	9.74	9.56	4.48	4.74	4.98	0.82	0.98	0.98
RV_{BDR}	8.86	9.34	9.88	4.50	4.88	5.18	0.80	1.00	0.98

Fonte: próprio autor.

Assim como o teste RV , o teste gradiente apresenta comportamento bastante liberal com taxas de rejeição nula bem próximas do teste RV . Por exemplo, na Tabela 24, quando $\mu_t \in (0.15, 0.83)$, $\alpha = 10\%$ e $n = 15$, os testes escore e RV apresentam taxas iguais a 22.88% e 25.08%, respectivamente.

O teste gradiente apresenta uma redução considerável na distorção de tamanho em todos os casos em que $n = 45$. Melhor comportamento que o mesmo apresenta é quando $\alpha = 1\%$ e $n = 45$.. Por exemplo, para o caso em que $\sigma^2 = 6.0$, $\mu_t \in (0.15, 0.83)$, $\alpha = 1\%$ e $n = 45$, o teste apresenta distorção de tamanho igual a 0.28%. Quando $n = 15$ e 30 as distorções são iguais a 4.18% e 0.96%, respectivamente.

Os testes baseados em reamostragem bootstrap mostraram excelente desempenho quanto ao tamanho. Por exemplo, na Tabela 23, quando $\mu_t \in (0.87, 0.98)$, $\alpha = 10\%$ e $n = 45$, os testes RV_{BB} e RV_{BDR} apresentaram taxas de rejeição iguais a 10.86% e 10.82%, respectiva-

mente.

No geral, assim como no modelo anterior, os tamanhos dos testes não foram muito afetados com os cenários considerados para o parâmetro de dispersão σ^2 . Portanto, o teste escore foi o que mais se destacou entre os testes assintóticos, pois foi o que apresentou menor distorção de tamanho. Nota-se, que os testes baseados em reamostragem bootstrap apresentaram desempenho superior aos demais testes. Podemos dizer que existem fortes indícios que esses testes tendem a controlar bem a probabilidade de erro tipo I, dado a habilidade em controlar bem as estimativas desta probabilidade.

Portanto, com base nos resultados apresentados, sugerimos usar testes baseados em reamostragem bootstrap para a realização de testes de hipóteses em pequenas amostras.

Tabela 21 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \gamma_2 = \gamma_3 = 0$ ($r = 2$). $\sigma^2 = 0.5$.

$\mu_t \in (0.01, 0.13)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	31.86	17.06	13.82	22.40	10.30	7.68	10.26	2.90	1.84
S_R	7.36	8.34	8.70	3.28	4.18	4.38	0.60	0.96	1.04
W	59.00	28.34	20.80	51.30	20.96	13.42	37.44	10.32	4.82
S_G	27.40	15.98	13.02	18.26	9.34	7.30	6.92	2.28	1.54
RV_{BB}	10.40	10.42	9.82	4.88	5.24	4.92	0.92	0.98	1.00
RV_{BDR}	9.62	10.24	10.04	5.08	5.08	4.82	0.82	0.96	1.00
$\mu_t \in (0.15, 0.83)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	30.00	16.44	13.78	20.78	9.74	7.36	9.20	2.84	1.80
S_R	8.84	9.82	9.56	4.18	5.08	4.88	1.00	1.02	1.00
W	52.80	26.36	19.24	45.32	18.08	12.30	33.96	8.42	4.00
S_G	26.50	15.44	13.22	17.24	8.76	6.80	6.60	2.18	1.40
RV_{BB}	10.68	10.36	9.98	5.60	5.30	4.86	1.08	1.18	0.90
RV_{BDR}	10.38	10.14	9.82	5.46	5.30	5.18	1.14	1.16	1.06
$\mu_t = (0.87, 0.98)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	28.26	16.14	14.86	19.22	9.02	8.18	8.04	2.74	2.02
S_R	7.80	8.88	10.32	3.90	4.30	4.96	0.76	1.08	1.04
W	49.86	26.18	20.46	41.72	17.80	13.62	28.86	8.22	4.70
S_G	24.14	15.16	14.28	15.36	7.84	7.60	5.56	2.18	1.78
RV_{BB}	9.94	9.84	11.16	5.10	5.14	5.44	0.98	1.00	1.20
RV_{BDR}	9.60	9.50	10.94	5.12	4.78	5.72	1.06	1.18	1.16

Fonte: próprio autor.

Tabela 22 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \gamma_2 = \gamma_3 = 0$ ($r = 2$). $\sigma^2 = 1.5$.

$\mu_t \in (0.01, 0.13)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	33.50	16.02	14.28	23.66	9.26	7.64	11.26	2.84	1.74
S_R	7.44	8.20	9.20	3.80	3.98	4.42	0.68	0.88	0.96
W	59.66	27.34	20.46	52.62	19.30	13.38	39.50	9.58	5.34
S_G	28.62	14.98	13.56	19.42	8.40	7.32	7.96	2.22	1.34
RV_{BB}	10.48	9.24	10.28	5.52	4.74	5.12	1.00	0.94	0.90
RV_{BDR}	10.00	9.14	10.12	5.40	4.80	5.26	1.22	0.90	1.06

$\mu_t \in (0.15, 0.83)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	28.50	15.76	13.02	19.36	8.68	6.64	7.88	2.18	1.60
S_R	8.44	9.14	9.32	4.16	4.66	4.68	1.02	1.10	1.04
W	51.90	25.40	18.94	43.44	17.22	11.92	30.76	7.50	3.76
S_G	24.64	14.72	12.46	16.36	7.86	6.24	5.72	1.66	1.48
RV_{BB}	9.72	10.08	9.82	4.88	4.60	4.62	1.04	0.94	0.92
RV_{BDR}	9.76	10.22	9.50	4.68	4.62	4.84	1.04	0.96	1.08

$\mu_t = (0.87, 0.98)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	28.26	16.30	13.80	19.88	9.46	7.40	8.24	2.52	1.88
S_R	7.80	10.14	9.28	3.98	4.80	4.86	0.94	1.08	0.80
W	50.32	24.92	19.10	42.12	16.80	12.18	29.66	7.84	4.58
S_G	24.20	15.14	12.86	15.78	8.52	6.96	5.42	1.82	1.52
RV_{BB}	9.98	10.20	9.82	4.90	4.98	5.12	0.88	0.98	1.00
RV_{BDR}	10.04	10.10	10.24	4.58	5.00	4.96	0.96	1.10	1.00

Fonte: próprio autor.

Tabela 23 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \gamma_2 = \gamma_3 = 0$ ($r = 2$). $\sigma^2 = 3.0$.

$\mu_t \in (0.01, 0.13)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	33.74	16.96	14.52	24.64	9.98	8.20	11.22	3.12	1.86
S_R	6.94	8.80	9.78	3.58	4.18	4.94	0.76	0.84	1.10
W	60.14	28.42	21.04	53.28	19.90	13.80	41.38	10.32	5.04
S_G	28.72	16.12	14.16	19.72	9.00	7.62	7.58	2.40	1.62
RV_{BB}	9.94	10.16	10.70	4.82	5.26	5.20	0.80	0.96	1.04
RV_{BDR}	9.68	9.96	10.28	4.68	5.18	5.28	0.84	1.14	1.24
$\mu_t \in (0.15, 0.83)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	28.04	15.84	12.88	19.02	9.32	7.02	8.30	2.24	1.58
S_R	7.90	9.04	9.08	3.80	4.74	4.64	0.88	1.06	1.02
W	50.08	24.96	18.52	42.14	17.54	11.44	29.72	7.94	3.78
S_G	25.20	15.06	12.44	16.58	8.60	6.50	6.36	1.68	1.28
RV_{BB}	10.46	10.76	9.80	5.36	5.16	4.78	1.22	0.70	0.88
RV_{BDR}	10.42	10.56	9.76	5.22	5.08	4.72	1.18	0.86	1.02
$\mu_t = (0.87, 0.98)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	28.70	17.00	14.36	18.74	9.40	8.14	8.14	2.94	1.90
S_R	8.18	9.36	10.42	3.80	4.58	5.06	0.94	0.98	1.00
W	50.50	26.16	19.24	43.22	18.28	12.14	30.46	8.24	4.64
S_G	24.20	16.02	13.78	15.16	8.36	7.36	5.16	2.22	1.52
RV_{BB}	9.56	10.20	10.86	4.58	5.32	5.44	0.94	1.16	0.98
RV_{BDR}	9.32	10.42	10.82	4.14	5.18	5.36	0.76	1.36	0.90

Fonte: próprio autor.

Tabela 24 – Taxas de rejeição nula (%) dos testes da razão de verossimilhanças generalizada (RV), escore (S_R), Wald (W), gradiente (S_G), razão de verossimilhanças generalizada corrigida via fator de Bartlett bootstrap (RV_{BB}) e da razão de verossimilhanças generalizada bootstrap duplo rápido (RV_{BDR}) em pequenas amostras. $H_0 : \gamma_2 = \gamma_3 = 0$ ($r = 2$). $\sigma^2 = 6.0$.

$\mu_t \in (0.01, 0.13)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	34.68	17.18	14.02	25.06	10.36	7.28	12.34	2.90	1.76
S_R	7.66	9.16	8.78	3.92	4.72	4.26	0.82	1.18	0.86
W	60.04	28.48	20.10	52.72	20.14	12.96	41.12	9.92	4.56
S_G	30.32	16.30	13.54	20.50	9.10	6.82	8.18	2.16	1.42
RV_{BB}	10.40	10.58	9.84	5.18	5.02	4.64	0.68	1.04	1.00
RV_{BDR}	10.08	10.34	9.74	5.08	5.10	4.84	0.90	1.02	0.84

$\mu_t \in (0.15, 0.83)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	25.08	15.08	12.50	16.66	8.78	6.80	6.86	2.24	1.44
S_R	7.20	9.12	8.94	3.34	4.42	4.48	0.74	0.90	1.00
W	45.82	23.40	18.40	37.74	16.38	11.24	26.24	7.14	3.74
S_G	22.88	14.20	11.90	14.12	7.88	6.56	5.18	1.96	1.28
RV_{BB}	9.28	10.38	9.96	4.84	5.04	4.80	1.24	0.98	0.96
RV_{BDR}	9.62	10.16	9.74	4.80	5.04	5.20	0.88	1.32	1.02

$\mu_t = (0.87, 0.98)$									
	$\alpha = 10\%$			$\alpha = 5\%$			$\alpha = 1\%$		
n	15	30	45	15	30	45	15	30	45
RV	28.98	15.48	14.04	20.12	8.68	7.24	8.24	2.92	1.56
S_R	8.20	9.00	9.92	3.96	4.26	5.12	0.86	1.14	0.86
W	50.90	23.70	19.54	43.64	16.32	12.04	31.74	7.48	4.32
S_G	25.56	14.36	13.48	15.52	8.04	6.86	5.84	1.98	1.42
RV_{BB}	9.82	9.30	9.98	4.94	5.06	5.12	1.04	1.18	0.80
RV_{BDR}	10.12	9.56	9.92	4.92	5.04	4.98	1.30	1.22	0.96

Fonte: próprio autor.

4.7 CRITÉRIOS PARA A SELEÇÃO DE MODELOS

Na inferência estatística é de fundamental importância escolher um modelo com menor número de parâmetros possíveis e que explique bem a aleatoriedade da variável resposta. Na literatura, existem diversos critérios para a seleção de modelos. É interessante que um modelo de regressão controle simultaneamente o viés do valor predito da t -ésima observação $\hat{\eta}_t$ e ao mesmo tempo a variância da resposta estimada pelo modelo de regressão ajustado a um conjunto de dados. Neste trabalho, faremos uso da estatística P^2 e sua versão corrigida P_c^2 , ambas propostas por Silva (2019) para modelos de regressão simplex não linear, nas quais avaliam o valor preditivo do modelo. Quanto à variância da resposta explicada pelo modelo de regressão ajustado vamos considerar as medidas R_{RV}^2 , R_{FC}^2 e suas respectivas versões corrigidas $R_{RV_c}^2$ e $R_{FC_c}^2$ propostas por Bayer e Cribari-Neto (2017) para o modelo de regressão beta.

A estatística *PRESS* que compõe o critério de predição P^2 se baseia na soma de quadrados dos resíduos. Neste sentido, vamos definir algumas quantidades relacionadas ao modelo de regressão simplex não linear.

Seja a matriz H^* dada por $H^* = (\widehat{W}\widehat{S})^{1/2}\widetilde{\mathcal{X}}(\widetilde{\mathcal{X}}\widehat{S}\widehat{W}\widetilde{\mathcal{X}})^{-1}\widetilde{\mathcal{X}}^\top(\widehat{S}\widehat{W})^{1/2}$, em que $\widetilde{\mathcal{X}} = \frac{\partial \eta}{\partial \beta}$ é uma matriz de derivadas $n \times k$, a matriz S está definida em (2.7) e W é uma matriz diagonal com elementos definidos em (2.11). Considere também a matriz diagonal U cujos elementos estão na Equação (2.6). A estatística *PRESS* para o modelo de regressão simplex não linear (SILVA, 2019) é definida por

$$PRESS = \sum_{t=1}^n (\check{y}_t - \hat{y}_{(t)})^2 = \sum_{t=1}^n \left(\frac{r_t}{1 - h_{tt}^*} \right)^2, \quad (4.3)$$

em que $\check{y}_t - \hat{y}_{(t)}$ é chamado erro predito com $\check{y} = \widehat{S}^{1/2}\widehat{W}^{1/2}u_1$ e $u_1 = \widehat{\eta} + \widehat{W}^{-1}\widehat{U}\widehat{T}(y - \widehat{\mu})$, h_{tt}^* é o t -ésimo elemento da diagonal principal da matriz H^* e r_{1t} é o resíduo ponderado proposto por Espinheira e Silva (2019) e definido por

$$r_{1t} = \frac{\hat{u}_t(y_t - \widehat{\mu}_t)}{\sqrt{\hat{b}_t}} \text{ com } b_t = \sigma_t^2 \left\{ \frac{3\sigma_t^2}{\mu_t(1 - \mu_t)} + \frac{1}{\mu_t^3(1 - \mu_t)^3} \right\}.$$

A estatística P^2 para o modelo de regressão simplex não linear (SILVA, 2019) é dada por

$$P^2 = 1 - \frac{PRESS}{SST_{(t)}^*}, \quad (4.4)$$

em que $SST_{(t)}^* = \sum_{t=1}^n (\check{y}_t - \bar{\check{y}}_{(t)})^2$ e $\bar{\check{y}}_{(t)}$ é a média aritmética dos $n - 1$ valores do vetor $\check{y} = \widehat{S}^{1/2}\widehat{W}^{1/2}u_1$ excluindo a t -ésima observação. Cox e Reid (1987) sugerem versões alternativas

da estatística *PRESS* baseada em diferentes resíduos. Assim, consideramos mais outros dois resíduos definidos abaixo e associamos as estatísticas *PRESS* e P^2 . O resíduo ponderado padronizado r_{2t} definido por Espinheira e Silva (2019) é dado por

$$r_{2t} = \frac{\hat{u}_t(y_t - \hat{\mu}_t)}{\sqrt{\hat{b}_t(1 - h_{tt}^*)}}.$$

Outro resíduo que utilizaremos é o resíduo combinado r_{3t} proposto por Silva (2019) e expresso como

$$r_{3t} = \frac{\hat{u}_t(y_t - \hat{\mu}_t) + \hat{a}_t}{\sqrt{\hat{b}_t + [2(\hat{\sigma}_t^2)^2]^{-1}}},$$

em que a_t é dado em (2.9).

Note que o valor numérico assumido pela estatística *PRESS* é positivo, assim a P^2 assume valores em $(-\infty; 1]$. Quanto mais próximo de um o valor de *PRESS* melhor poder preditivo do modelo. Consideramos também a versão corrigida da estatística P^2 apresentada em Silva (2019) e dada por $P_c^2 = 1 - \frac{(1-P^2)(n-1)}{(n-(k_1+q_1))}$, em que k_1 e q_1 são os números de covariáveis do submodelo da média e dispersão, respectivamente. As estatísticas P^2 associadas aos resíduos r_{1t} , r_{2t} e r_{3t} , serão denotadas por P_1^2 , P_2^2 e P_3^2 , respectivamente, e P_{1c}^2 , P_{2c}^2 e P_{3c}^2 são suas respectivas versões corrigidas.

Para avaliar a qualidade de ajuste do modelo aos dados originais consideramos o R_{RV}^2 (NAGELKERKE, 1991) e sua versão corrigida $R_{RV_c}^2$ (BAYER; CRIBARI-NETO, 2017) dados, respectivamente, por $R_{RV}^2 = 1 - \left(\frac{L_{null}}{L_{fit}}\right)^{2/n}$ e $R_{RV_c}^2 = 1 - (1 - R_{RV}^2) \left(\frac{n-1}{n-(1+\alpha)k_1-(1-\alpha)q_1}\right)^\delta$, em que L_{null} é a função de verossimilhança maximizada do modelo sem regressores, L_{fit} a função de verossimilhança maximizada do modelo ajustado, $0 \leq \alpha \leq 1$ e $\delta > 0$. Bayer e Cribari-Neto (2017) sugerem $\alpha = 0.4$ e $\delta = 1$.

Consideramos ainda a medida R_{FC}^2 definida por Ferrari e Cribari-Neto (2004) como o quadrado da correlação entre $g(y)$ e $\hat{\eta}$. Baseado nela, Bayer e Cribari-Neto (2017) definiram a versão corrigida $R_{FC_c}^2$ para o modelo de regressão beta dada por $R_{FC_c}^2 = \frac{1-(1-R_{FC}^2)(n-1)}{(n-(k_1+q_1))}$.

4.8 APLICAÇÃO: DADOS DE CRAQUEAMENTO CATALÍTICO FLUIDO (FCC)

Nesta seção, será apresentada uma aplicação dos testes anteriormente avaliados ao conjunto de dados modelados por Espinheira e Silva (2019), a respeito do processo de craqueamento catalítico fluido. Foi considerado o modelo de regressão simplex não linear definido por

$$M1 : \Phi^{-1}(\mu_t) = \beta_1 + \beta_2 \frac{x_{t2}}{x_{t2} + \beta_3} + \beta_4 x_{t3} + \beta_5 \sqrt{x_{t4}},$$

$$\log(\sigma_t^2) = \gamma_1 + \gamma_2 x_{t4}^2,$$

com $t = 1, \dots, 28$ e $x_{t2} + \beta_3 \neq 0$ para todo t .

Inicialmente o modelo foi avaliado com relação ao fator de não linearidade. Vale notar que a hipótese nula $H_0 : \beta_2 = 0$ leva a problemas de identificabilidade e a uma matriz de informação singular sob H_0 . De fato, a t -ésima linha da matriz $\tilde{\mathcal{X}}$ é definida por $(1, x_{t2}/(x_{t2} + \beta_3), -\beta_2 x_{t2}/(x_{t2} + \beta_3)^{-2}, x_{t3}, \sqrt{x_{t4}})$. Sob $H_0 : \beta_2 = 0$, tal matriz é singular o que implica matriz de informação singular. Em tais condições, os testes usuais não podem ser utilizados, já que as condições de regularidade (ver Apêndice B) que garantem a aproximação assintótica dos mesmos não são respeitadas (SILVEY, 1959; GALLANT, 1977; BARNABANI, 2006). Para contornar tal situação utilizamos o fato que $\beta_3 \neq 0$, já que x_{t2} assume valores nulos, e substituímos o mesmo por sua estimativa antes de realizar o teste sobre β_2 . A Tabela 25 apresenta as estimativas dos parâmetros e erros-padrão para o modelo $M1$.

Tabela 25 – Estimativas de máxima verossimilhança $\hat{\theta}$ e erros-padrão (EP) para os parâmetros do modelo $M1$. Dados de craqueamento catalítico fluido (FCC).

Parâmetros	β_1	β_2	β_3	β_4	β_5	γ_1	γ_2
$\hat{\theta}$	1.3994	-0.0609	-27.8429	-0.1820	-0.4201	0.7561	-1.1502
EP($\hat{\theta}$)	0.0769	0.0227	3.8788	0.0629	0.0747	0.3698	0.3143

Fonte: próprio autor.

A Tabela 26 apresenta os p -valores obtidos para o teste da hipótese nula $H_0 : \beta_2 = 0$ para o modelo

$$M2 : \Phi^{-1}(\mu_t) = \beta_1 + \beta_2 w_{t2} + \beta_4 x_{t3} + \beta_5 \sqrt{x_{t4}},$$

$$\log(\sigma_t^2) = \gamma_1 + \gamma_2 x_{t4}^2,$$

com $w_{t2} = x_{t2}/(x_{t2} + \hat{\beta}_3)$ e $t = 1, \dots, 28$. Todos os testes rejeitam a hipótese nula ao nível nominal de 5%.

Em seguida, foi considerada a hipótese nula $H_0 : \gamma_1 = \gamma_2 = 0$, que corresponde a testar o aspecto de dispersão variável no modelo. A Tabela 27 contém os p -valores para tal teste. Os

Tabela 26 – p -valores dos testes de $H_0 : \beta_2 = 0$.

	RV	S_R	W	S_G	RV_{BB}	RV_{BDR}
p -valor	0.0007	0.0066	0.0000	0.0036	0.0047	0.0160

Fonte: próprio autor.

Tabela 27 – p -valores dos testes de $H_0 : \gamma_1 = \gamma_2 = 0$.

	RV	S_R	W	S_G	RV_{BB}	RV_{BDR}
p -valor	0.0187	0.0843	0.0010	0.0210	0.0030	0.0200

Fonte: próprio autor.

resultados indicam que ao nível nominal de 5% apenas o teste escore indica a não rejeição da hipótese nula. Logo, seguimos com o modelo considerando dispersão variável.

Dando continuidade aos testes, foram realizados testes individuais sobre os parâmetros dos fatores lineares do submodelo da média e os parâmetros do submodelo da dispersão. A Tabela 28 apresenta os p -valores obtidos para tais testes. Os resultados mostram que ao nível nominal de 5% os testes não indicam a exclusão de nenhuma covariável do submodelo da média. Por sua vez, com relação ao submodelo da dispersão, os resultados mostram que ao nível nominal de 5% os testes assintóticos não indicam a exclusão de nenhuma covariável. No entanto, os p -valores dos testes baseados em bootstrap, RV_{BB} e RV_{BDR} , indicam a não rejeição da hipótese nula $\gamma_1 = 0$ mesmo ao nível nominal de 10%.

Tabela 28 – p -valores dos testes individuais sobre os parâmetros do modelo $M2$.

H_0	RV	S_R	W	S_G	RV_{BB}	RV_{BDR}
$\beta_1 = 0$	< 0.0001	< 0.0001	< 0.0001	< 0.0001	< 0.0001	< 0.0001
$\beta_4 = 0$	0.0082	0.0160	0.0038	0.0116	0.0188	0.0160
$\beta_5 = 0$	< 0.0001	0.0055	0.0000	0.0022	0.0019	0.0060
$\gamma_1 = 0$	0.0377	0.0189	0.0409	0.0284	0.1061	0.1240
$\gamma_2 = 0$	0.0048	0.0491	0.0002	0.0073	0.0150	0.0200

Fonte: próprio autor.

Logo, nosso modelo final é definido por

$$M3 : \Phi^{-1}(\mu_t) = \beta_1 + \beta_2 \frac{x_{t2}}{x_{t2} + \beta_3} + \beta_4 x_{t3} + \beta_5 \sqrt{x_{t4}},$$

$$\log(\sigma_t^2) = \gamma_2 x_{t4}^2,$$

com $t = 1, \dots, 28$.

A Tabela 29 contém critérios de seleção para os modelos $M2$ e $M3$. Os resultados indicam a superioridade do modelo $M3$ para a maior parte das medidas adotadas.

Tabela 29 – Critérios de seleção para os modelos $M2$ e $M3$.

	P_1^2	P_{1c}^2	P_2^2	P_{2c}^2	P_3^2	P_{3c}^2	R_{FC}^2	R_{FCc}^2	R_{RV}^2	R_{RVc}^2
$M2$	0.7274	0.6495	0.6560	0.5577	0.7860	0.7249	0.6506	0.5508	0.7228	0.6435
$M3$	0.8145	0.7723	0.7733	0.7217	0.8231	0.7830	0.6605	0.5833	0.6765	0.6030

Fonte: próprio autor.

Por sua vez, a Tabela 30 apresenta as estimativas de máxima verossimilhança para os parâmetros do modelo selecionado e seus respectivos erros-padrão.

Tabela 30 – Estimativas de máxima verossimilhança $\hat{\theta}$ e erros-padrão (EP) para os parâmetros do modelo $M3$. Dados de craqueamento catalítico fluido (FCC).

Parâmetros	β_1	β_2	β_3	β_4	β_5	γ_2
$\hat{\theta}$	1.4112	-0.0626	-26.7268	-0.2401	-0.4086	-0.5886
EP($\hat{\theta}$)	0.0624	0.0161	3.5272	0.0606	0.0599	0.2271

Fonte: próprio autor.

Como vimos anteriormente, a utilização dos testes baseados em bootstrap RV_{BB} e RV_{BDR} foram cruciais para chegarmos no modelo final $M3$. Os p -valores apresentados por eles para o parâmetro $\gamma_1 = 0$ foram 0.1061 e 0.1240, respectivamente. Indicando a não rejeição da hipótese nula $H_0 : \gamma_1 = 0$ mesmo ao nível nominal de 10%. No estudo de simulação de Monte Carlo foi observado que os testes via bootstrap apresentaram superioridade quando comparado aos testes assintóticos, reforçando a importância da utilização desse método. Portanto, com base nos resultados apresentados sugere-se a utilização dos testes corrigidos em pequenas amostras.

5 CONSIDERAÇÕES FINAIS

Neste trabalho propomos melhoramentos inferenciais baseados em bootstrap para os parâmetros que indexam o modelo de regressão simplex não linear (ESPINHEIRA; SILVA, 2019) em pequenas amostras. Muitas vezes os EMVs podem ser bastante viesados quando o tamanho da amostra é pequeno ou até mesmo moderado. Uma das formas de contornar esse problema é usando o método bootstrap, pois através do mesmo é possível obter estimadores menos viesados em amostras finitas. Assim, foi considerado o estimador de máxima verossimilhança e uma versão do EMV corrigida via esquema bootstrap. Para avaliar a qualidade dos estimadores consideramos o viés relativo e a raiz do erro quadrático médio. Os resultados de simulações de Monte Carlo apresentados no terceiro capítulo mostraram que de modo geral, os estimadores corrigidos apresentaram vieses inferiores aos dos estimadores de máxima verossimilhança, evidenciando a efetividade do esquema bootstrap na correção do viés. Percebemos também que as estimativas da $\sqrt{EQM}$ de todos os estimadores diminuíram ao aumentarmos o tamanho da amostra. De modo geral, os estimadores de máxima verossimilhança e os corrigidos via esquema bootstrap apresentaram bom desempenho em pequenas amostras.

Conforme vimos no terceiro capítulo, os intervalos de confiança assintóticos requerem grandes amostras para que as coberturas exatas estejam próximas às nominais. Uma alternativa para a construção de intervalos de confiança adequados em pequenas amostras é através do método bootstrap. Neste trabalho, fizemos uso dos intervalos de confiança bootstrap percentil e bootstrap- t . Avaliamos seus respectivos desempenhos e comparamos com os resultados obtidos com base nos intervalos de confiança assintóticos. Percebemos de modo geral que os intervalos de confiança assintóticos e bootstrap percentil não obtiveram resultados tão satisfatórios como o do intervalo de confiança bootstrap- t . O ICBT apresentou excelentes taxas de cobertura para os parâmetros do modelo de regressão simplex não linear. Adicionalmente, foi feito um estudo para os parâmetros de não linearidade dos submodelos da média e dispersão, β_2 e γ_2 , respectivamente. Avaliamos diversos aspectos da estimação intervalar, a saber: o limite inferior (superior) médio, o comprimento médio, a probabilidade empírica de cobertura e os percentuais de vezes que o verdadeiro valor do parâmetro foi menor (maior) que o limite inferior (superior) do IC. Com base nos resultados de simulações sugere-se o uso do intervalo de confiança bootstrap- t que apresentou melhor cobertura e balanceamento.

No quarto capítulo abordamos testes de hipóteses. Primeiramente, foram considerados os

testes da razão de verossimilhanças generalizada, escore, Wald e gradiente. Esses testes são realizados a partir de valores críticos válidos para grandes amostras, o que pode causar considerável distorção no tamanho do teste em amostras pequenas ou moderadas. Essa distorção pode ser contornada através de correções entre as quais baseadas na técnica de reamostragem bootstrap. Consideramos duas versões bootstrap baseadas no teste da razão de verossimilhanças generalizada, a versão corrigida via fator de Bartlett bootstrap e o teste bootstrap duplo rápido. Avaliamos os seus respectivos desempenhos e consideramos testes sobre os parâmetros que indexam os submodelos da média e dispersão. Percebemos que no geral, dentre os testes assintóticos considerados, o escore mostrou melhor desempenho com taxas de rejeição nula mais próximas do nível nominal de referência. Os testes baseados em reamostragem bootstrap apresentaram desempenho superior aos demais testes, mostrando ser eficaz no controle da probabilidade de erro tipo I. Portanto, com base nos resultados apresentados, sugerimos usar testes baseados em reamostragem bootstrap para a realização de testes de hipóteses em pequenas amostras.

5.1 TRABALHOS FUTUROS

Serão objetivos de nossas próximas pesquisas:

- Implementação de intervalos de predição bootstrap para o modelo de regressão simplex linear e não linear;
- Considerar outros modelos de regressão simplex não linear e avaliar os desempenhos dos testes nos mesmos.

REFERÊNCIAS

- ATKINSON, A. C. *Plots, transformations and regression: an introduction to graphical methods of diagnostic regression analysis*. New York: Oxford University Press, 1985.
- BARNABANI, M. Maximum likelihood estimator and singularity of the information matrix. *Quaderni di Statistica*, v. 8, 2006.
- BARNDORFF-NIELSEN, O. E.; JØRGENSEN, B. Some parametric models on the simplex. *Journal of Multivariate Analysis*, v. 39, n. 1, p. 106–116, 1991.
- BARTLETT, M. S. Properties of sufficiency and statistical tests. *Proceedings of the Royal Society of London. Series A-Mathematical and Physical Sciences*, v. 160, n. 901, p. 268–282, 1937.
- BAYER, F. M.; CRIBARI-NETO, F. Bartlett corrections in beta regression models. *Journal of Statistical Planning and Inference*, v. 143, n. 3, p. 531–547, 2013.
- BAYER, F. M.; CRIBARI-NETO, F. Model selection criteria in beta regression with varying dispersion. *Communications in Statistics-Simulation and Computation*, v. 46, n. 1, p. 729–746, 2017.
- BUSE, A. The likelihood ratio, wald, and lagrange multiplier tests: an expository note. *The American Statistician*, v. 36, n. 3a, p. 153–157, 1982.
- COX, D. R.; HINKLEY, D. V. *Theoretical statistics*. London: Chapman and Hall, 1974.
- COX, D. R.; REID, N. Parameter orthogonality and approximate conditional inference. *Journal of the Royal Statistical Society: Series B (Methodological)*, v. 49, n. 1, p. 1–18, 1987.
- CRIBARI-NETO, F.; LIMA, F. P. Resampling-based prediction intervals in beta regressions under correct and incorrect model specification. *Communications in Statistics-Simulation and Computation*, p. 1–19, 2019.
- CRIBARI-NETO, F.; ZARKOS, S. G. R: Yet another econometric programming environment. *Journal of Applied Econometrics*, v. 14, n. 3, p. 319–329, 1999.
- DALGAARD, P. *Introductory statistics with R*. New York: Springer, 2002.
- DAVIDSON, R.; MACKINNON, J. G. Improving the reliability of bootstrap tests. Working Papers 995, Queen's University, Department of Economics, 2000.
- DAVIDSON, R.; MACKINNON, J. G. Fast double bootstrap tests of nonnested linear regression models. *Econometric Reviews*, v. 21, n. 4, p. 419–429, 2002.
- DAVIDSON, R.; MACKINNON, J. G. Improving the reliability of bootstrap tests with the fast double bootstrap. *Computational Statistics & Data Analysis*, v. 51, n. 7, p. 3259–3281, 2007.
- DAVISON, A. C.; HINKLEY, D. V. *Bootstrap methods and their application*. New York: Cambridge University Press, 1997.
- DOORNIK, J. A. *Ox: an object-oriented matrix programming language*. London: Timberlake Consultants LTD, 2001.

- DOORNIK, J. A. *Object-oriented matrix programming using Ox*. London: Timberlake Consultants Press, 2006.
- DOORNIK, J. A. *An object-oriented matrix language Ox 6*. London: Timberlake Consultants Press, 2009.
- EFRON, B. Bootstrap methods: another look at the jackknife. *Annals of Statistics*, v. 7, n. 1, p. 1–26, 1979.
- EFRON, B. Estimating the error rate of a prediction rule: improvement on cross-validation. *Journal of the American Statistical Association*, v. 78, n. 382, p. 316–331, 1983.
- EFRON, B.; TIBSHIRANI, R. J. *An introduction to the bootstrap*. New York: Chapman & Hall, 1994.
- ESPINHEIRA, P. L.; SILVA, A. Residual and influence analysis to a general class of simplex regression. *Test*, p. 1–30, 2019.
- FERRARI, S.; CRIBARI-NETO, F. Beta regression for modelling rates and proportions. *Journal of Applied Statistics*, v. 31, n. 7, p. 799–815, 2004.
- FERRARI, S. L.; CRIBARI-NETO, F. On bootstrap and analytical bias corrections. *Economics Letters*, v. 58, n. 1, p. 7–15, 1998.
- FERRARI, S. L. P.; PINHEIRO, E. C. Improved likelihood inference in beta regression. *Journal of Statistical Computation and Simulation*, v. 81, n. 4, p. 431–443, 2011.
- GALLANT, A. R. Testing a nonlinear regression specification: a nonregular case. *Journal of the American Statistical Association*, v. 72, n. 359, p. 523–530, 1977.
- GUEDES, A. C.; CRIBARI-NETO, F.; ESPINHEIRA, P. L. Modified likelihood ratio tests for unit gamma regressions. *Journal of Applied Statistics*, p. 1–25, 2020.
- JØRGENSEN, B. *The theory of dispersion models*. London: Chapman and Hall, 1997.
- KNUTH, D. *The TEXbook*. New York: Tadisson-Wesley, 1986.
- LAMPORT, L. *LATEX: a document preparation system: user's guide and reference manual*. Massachusetts: Addison-Wesley, 1994.
- LEMONTE, A. *The gradient test: another likelihood-based test*. London: Academic Press, 2016.
- LEMONTE, A. J.; BAZAN, J. L. New class of Johnson S_B distributions and its associated regression model for rates and proportions. *Biometrical Journal*, v. 58, p. 727–746, 2016.
- LIMA, F. P.; CRIBARI-NETO, F. Bootstrap-based testing inference in beta regressions. *Brazilian Journal of Probability and Statistics*, v. 34, n. 1, p. 18–34, 2020.
- LIU, P.; YUEN, K. C.; WU, L.-C.; TIAN, G.-L.; LI, T. Zero-one-inflated simplex regression models for the analysis of continuous proportion data. *Statistics and Its Interface*, v. 13, n. 2, p. 193–208, 2020.
- LÓPEZ, F. O. A bayesian approach to parameter estimation in simplex regression model: a comparison with beta regression. *Revista Colombiana de Estadística*, v. 36, n. 1, p. 1–21, 2013.

- MACKINNON, J. G. Applications of the fast double bootstrap. Kingston (Ontario): Queen's University, Department of Economics, 2006.
- MCCULLAGH, P.; NELDER, J. A. *Generalized linear models*. London: Chapman and Hall, 1989.
- MITNIK, P. A.; BAEK, S. The Kumaraswamy distribution: median-dispersion re-parameterizations for regression modeling and simulation-based estimation. *Statistical Papers*, v. 54, n. 1, p. 177–192, 2013.
- MIYASHIRO, E. S. *Modelos de regressão beta e simplex para análise de proporções*. Tese (Doutorado) — Universidade de São Paulo, 2008.
- MOUSA, A. M.; EL-SHEIKH, A. A.; ABDEL-FATTAH, M. A. A gamma regression for bounded continuous variables. *Advances and Applications in Statistics*, v. 49, n. 4, p. 305–326, 2016.
- MUGGEO, V. M.; LOVISON, G. The “three plus one” likelihood-based test statistics: unified geometrical and graphical interpretations. *The American Statistician*, v. 68, n. 4, p. 302–306, 2014.
- NAGELKERKE, N. J. A note on a general definition of the coefficient of determination. *Biometrika*, v. 78, n. 3, p. 691–692, 1991.
- NELDER, J. A.; WEDDERBURN, R. W. Generalized linear models. *Journal of the Royal Statistical Society: Series A (General)*, v. 135, n. 3, p. 370–384, 1972.
- OSPINA, R.; CRIBARI-NETO, F.; VASCONCELLOS, K. L. Improved point and interval estimation for a beta regression model. *Computational Statistics & Data Analysis*, v. 51, n. 2, p. 960–981, 2006.
- PAWITAN, Y. *In all likelihood: statistical modelling and inference using likelihood*. United Kingdom: Oxford University Press, 2001.
- RAO, C. R. Large sample tests of statistical hypotheses concerning several parameters with applications to problems of estimation. *Proceedings of the Cambridge Philosophical Society*, v. 44, n. 1, p. 50–57, 1948.
- ROCHA, S. S.; ESPINHEIRA, P. L.; CRIBARI-NETO, F. Residual and local influence analyses for unit gamma regressions. *Statistica Neerlandica*, v. 75, n. 2, p. 137–160, 2021.
- ROCKE, D. M. Bootstrap bartlett adjustment in seemingly unrelated regression. *Journal of the American Statistical Association*, v. 84, n. 406, p. 598–601, 1989.
- SALAZAR, S. M. G. *Contribución al estudio de la reacción de decomposición de la Zeolita Y en presencia de vapor de agua y vanadio*. Dissertação (Mestrado) — Universidad Nacional de Colombia, 2005.
- SANTOS, L. A. d. *Modelos de regressão simplex: resíduos de Pearson corrigidos e aplicações*. Tese (Doutorado) — Universidade de São Paulo, 2011.
- SILVA, F. C. d. *Teste de diagnóstico baseado em influência local aplicado ao modelo de regressão simplex*. Dissertação (Mestrado) — Universidade Federal de Pernambuco, 2016.

SILVA, L. C. M. d. *Coeficientes de predição para os modelos de regressão beta e simplex*. Dissertação (Mestrado) — Universidade Federal de Pernambuco, 2015.

SILVA, L. C. M. d. *Diagnóstico em modelos de regressão simplex*. Tese (Doutorado) — Universidade Federal de Pernambuco, 2019.

SILVEY, S. D. The lagrangian multiplier test. *The Annals of Mathematical Statistics*, v. 30, n. 2, p. 389–407, 1959.

SIMAS, A. B.; BARRETO-SOUZA, W.; ROCHA, A. V. Improved estimators for a general class of beta regression models. *Computational Statistics & Data Analysis*, v. 54, n. 2, p. 348–366, 2010.

SKOVGAARD, I. M. Likelihood asymptotics. *Scandinavian Journal of Statistics*, v. 28, n. 1, p. 3–32, 2001.

SMITHSON, M.; SHOU, Y. CDF-quantile distributions for modelling random variables on the unit interval. *British Journal of Mathematical and Statistical Psychology*, v. 70, p. 412–438, 2017.

SONG, P. X.-K. Dispersion models in regression analysis. *Pakistan Journal of Statistics*, v. 25, n. 4, 2009.

SONG, P. X.-K.; QIU, Z.; TAN, M. Modelling heterogeneous dispersion in marginal models for longitudinal proportional data. *Biometrical Journal: Journal of Mathematical Methods in Biosciences*, v. 46, n. 5, p. 540–553, 2004.

SONG, P. X.-K.; TAN, M. Marginal models for longitudinal continuous proportional data. *Biometrics*, v. 56, n. 2, p. 496–502, 2000.

TERRELL, G. R. The gradient statistic. *Computing Science and Statistics*, v. 34, n. 34, p. 206–215, 2002.

VARGAS, T. M.; FERRARI, S. L.; LEMONTE, A. J. Gradient statistic: higher-order asymptotics and bartlett-type correction. *Electronic Journal of Statistics*, v. 7, p. 43–61, 2013.

VENABLES, W.; RIPLEY, B. *Modern applied statistics with S*. New York: Springer, 2002.

VENEZUELA, M. K. *Equação de estimação generalizada e influência local para modelos de regressão beta com medidas repetidas*. Tese (Doutorado) — Universidade de São Paulo, 2008.

WALD, A. Tests of statistical hypotheses concerning several parameters when the number of observations is large. *Transactions of the American Mathematical Society*, v. 54, n. 3, p. 426–482, 1943.

WILKS, S. S. The large-sample distribution of the likelihood ratio for testing composite hypotheses. *The Annals of Mathematical Statistics*, v. 9, n. 1, p. 60–62, 1938.

APÊNDICE A – FUNÇÃO ESCORE PARA O PARÂMETRO μ

Seja y uma variável aleatória seguindo a distribuição simplex, $y \sim S^-(\mu, \sigma^2)$, o logaritmo da função densidade da distribuição simplex é dado por

$$\ell(\mu, \sigma^2) = -\frac{1}{2} \log 2\pi - \frac{1}{2} \log \sigma^2 - \frac{3}{2} \log\{y(1-y)\} - \frac{1}{2\sigma^2} d(y; \mu).$$

Podemos verificar que a função escore para o parâmetro μ é dada por

$$\begin{aligned} \frac{\partial \ell(\mu, \sigma^2)}{\partial \mu} &= \frac{\partial}{\partial \mu} \left\{ -\frac{1}{2\sigma^2} d(y; \mu) \right\} \\ &= -\frac{1}{2\sigma^2} \frac{\partial d(y; \mu)}{\partial \mu} \\ &= -\frac{1}{2\sigma^2} d'(y; \mu), \end{aligned}$$

em que,

$$\begin{aligned} d'(y; \mu) &= \frac{\partial}{\partial \mu} \left[\frac{(y - \mu)^2}{y(1-y)\mu^2(1-\mu)^2} \right] \\ &= \frac{2(y - \mu)(-1)[y(1-y)\mu^2(1-\mu)^2] - (y - \mu)^2 y(1-y)[2\mu(1-\mu)^2 - 2(1-\mu)\mu^2]}{y^2(1-y)^2\mu^4(1-\mu)^4} \\ &= \frac{-2(y - \mu)y(1-y)[\mu^2(1-\mu)^2 + (y - \mu)\mu(1-\mu)^2 - (y - \mu)\mu^2(1-\mu)]}{y^2(1-y)^2\mu^4(1-\mu)^4} \\ &= \frac{-2(y - \mu)y(1-y)\mu(1-\mu)[\mu(1-\mu) + (y - \mu)(1-\mu) - (y - \mu)\mu]}{y^2(1-y)^2\mu^4(1-\mu)^4} \\ &= \frac{-2(y - \mu)[\mu(1-\mu) + (y - \mu)(1-\mu) - (y - \mu)\mu]}{y(1-y)\mu^3(1-\mu)^3} \\ &= \frac{-2(y - \mu)[\mu - \mu^2 + (y - y\mu - \mu + \mu^2) - y\mu + \mu^2]}{y(1-y)\mu^3(1-\mu)^3} \\ &= \frac{-2(y - \mu)[y - 2y\mu + \mu^2]}{y(1-y)\mu^3(1-\mu)^3} \\ &= \frac{-2(y - \mu)[\mu^2 - 2y\mu + y + y^2 - y^2]}{y(1-y)\mu^3(1-\mu)^3} \\ &= \frac{-2(y - \mu)}{\mu(1-\mu)} \left[\frac{(y^2 - 2y\mu + \mu^2) + (y - y^2)}{y(1-y)\mu^2(1-\mu)^2} \right] \\ &= \frac{-2(y - \mu)}{\mu(1-\mu)} \left[\frac{(y - \mu)^2 + y(1-y)}{y(1-y)\mu^2(1-\mu)^2} \right] \\ &= \frac{-2(y - \mu)}{\mu(1-\mu)} \left[\frac{(y - \mu)^2}{y(1-y)\mu^2(1-\mu)^2} + \frac{y(1-y)}{y(1-y)\mu^2(1-\mu)^2} \right] \\ &= \frac{-2(y - \mu)}{\mu(1-\mu)} \left[d(y; \mu) + \frac{1}{\mu^2(1-\mu)^2} \right]. \end{aligned}$$

Logo, a função escore de μ pode ser expressa por

$$\frac{\partial \ell(\mu, \sigma^2)}{\partial \mu} = \frac{1}{\sigma^2} \frac{(y - \mu)}{\mu(1-\mu)} \left[d(y; \mu) + \frac{1}{\mu^2(1-\mu)^2} \right].$$

APÊNDICE B – CONDIÇÕES DE REGULARIDADE (C.R.)

As condições seguintes de regularidade são usadas na teoria assintótica para justificar e delimitar os erros de aproximação em expansões de séries de Taylor. Algumas dessas condições ou a totalidade delas são necessárias para provar propriedades como consistência, unicidade, eficiência e suficiência dos estimadores de máxima verossimilhança. Suponha que os dados y_i 's são realizações independentes identicamente distribuídas de uma variável aleatória y caracterizada por distribuições P_θ pertencentes a uma certa classe G , que dependem de um vetor θ de dimensão p , $\theta \in \Theta$. Sejam $f(y; \theta)$ e $L(\theta) = \prod_{i=1}^n f(y_i; \theta)$ as funções de probabilidade ou densidade comum dos dados e de verossimilhança para θ , respectivamente. Assim, consideramos as seguintes suposições (COX; HINKLEY, 1974), capítulo 9:

(i) As distribuições P_θ são identificáveis, isto é, $\theta \neq \theta'$ implica $P_\theta \neq P_{\theta'}$. Essa condição assegura que as distribuições de probabilidade dos dados definidas por dois valores distintos de θ são diferentes.

(ii) As distribuições P_θ têm o mesmo suporte para todo $\theta \in \Theta$, ou seja, o conjunto $A = \{y : f(y; \theta) > 0\}$ independe de θ (o suporte da distribuição de y não depende de θ). Essa condição garante que seus campos de variação são idênticos e independem de θ .

As suposições (iii) – (v) abaixo, garantem a regularidade de $f(y; \theta)$ como função de θ e a existência de um conjunto aberto Θ_1 no espaço paramétrico tal que o parâmetro θ_0 pertença a Θ_1 :

(iii) Existe um conjunto aberto Θ_1 em Θ contendo θ_0 tal que a função densidade $f(y; \theta)$ admite derivadas de terceira ordem em relação a θ para todo $\theta \in \Theta_1$. Essa condição garante a existência das derivadas até a terceira ordem de $f(y; \theta)$ em Θ_1 .

(iv) $\mathbb{E}[U(\theta)] = 0$ e a matriz $K(\theta) = \mathbb{E}[U(\theta)U(\theta)^\top]$ existe e é positiva-definida para todo $\theta \in \Theta_1$, em que $U(\theta)$ é a função score e $K(\theta)$ a matriz de informação de Fisher de θ . Essa condição diz que a matriz $K(\theta)$ é finita e positiva-definida numa vizinhança aberta de θ_0 , o que é uma condição suficiente para garantir que $\hat{\theta}$ é um máximo local nesta vizinhança.

(v) Existem funções $M_{ijk}(y)$ não envolvendo θ tais que, para $i, j, k = 1, \dots, p$,

$$\left| \frac{\partial^3 \log f(y; \theta)}{\partial \theta_i \partial \theta_j \partial \theta_k} \right| < M_{ijk}(y),$$

para todo $\theta \in \Theta_1$, em que $\mathbb{E} [M_{ijk}(y)] < \infty$. Essa condição diz que as terceiras derivadas do logaritmo da função de verossimilhança são limitadas por uma função integrável de y cuja esperança é finita.