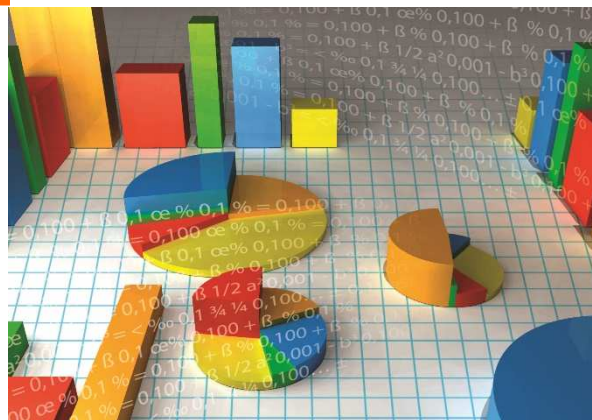


REVISTA BRASILEIRA DE ESTATÍSTICA

ISSN 2675-3243



volume 78

número 243

janeiro/junho 2020



IBGE
Instituto Brasileiro de Geografia e Estatística

Instituto Brasileiro de Geografia e Estatística – IBGE

REVISTA BRASILEIRA DE ESTATÍSTICA

Volume 78 - número 243 – jan/jun 2020

ISSN 2675-3243

R. Bras. Estat., Rio de Janeiro, v.78, n.243, p. 1-111, jan/jun 2020

Instituto Brasileiro de Geografia e Estatística – IBGE

@IBGE. 2020

Revista Brasileira de Estatística, ISSN 2675-3243

Órgão oficial do IBGE e da Associação Brasileira de Estatística – ABE. Publicação semestral que se destina a promover e ampliar o uso de métodos estatísticos através de divulgação de artigos inéditos tratando de aplicações da Estatística nas mais diversas áreas do conhecimento. Temas abordando aspectos do desenvolvimento metodológico serão aceitos, desde que relevantes para a produção e uso de estatísticas públicas.

Os originais para publicação deverão ser submetidos para o site <http://rbes.net.br>.

Os artigos submetidos à RBE não devem ter sido publicados em outros periódicos.

A Revista não se responsabiliza pelos conceitos emitidos em matéria assinada.

Editor Responsável

José André de Moura Brito - ENCE/IBGE

Editores Executivos

Marcel de Toledo Vieira - UFJF

Paulo Justiniano Ribeiro Junior – UFPR

Editores Associados

Alinne de Carvalho Veiga, ENCE/IBGE

Cristiano Ferraz, UFPE

Denise Britz do Nascimento Silva, ENCE/IBGE

Fernando Antônio Basile Colugnati, UFJF

Francisco Louzada-Neto, USP

Gustavo da Silva Ferreira, ENCE/IBGE

Ismenia Blavatsky, UFRN

Juvencio Santos Nobre, UFC

Maysa Sacramento de Magalhães, ENCE/IBGE

Paulo de Martino Jannuzzi, ENCE/IBGE

Pedro Luis do Nascimento Silva, ENCE/IBGE

Taiane Schaedler Prass, UFRGS

Vera Lucia Damasceno Tomazella, UFSCAR

Viviana Giampaoli, USP

Editoração

Luciane Canuto de Souza

Tamires de Assis Silva Sampaio

Capa

Ilustração da Capa

Informações Adicionais

Revista Brasileira de Estatística/IBGE - v.1, n.1 (jan/mar. 1940) – Rio de Janeiro: IBGE, 1940.v. trimestral (1940-1986), semestral (1987-).

Continuação de: Revista de economia e estatística. Índices acumulados de autor e assunto publicados no v.43 (1940-1979) e v.50 (1980-1989). Co-edição com a Associação Brasileira de Estatística a partir do v.58.

ISSN publicação online 2675-3243, a partir de 2019.

I. Estatística – Periódicos. II. IBGE. III. Associação Brasileira de Estatística.

Gerência de Biblioteca e Acervos Especiais CDU 31(05)
RJ-IBGE/88-05 (ver. 2009) PERIÓDICO

Sumário

Nota do Editor	5
----------------------	---

Artigos

A nálise da lei de benford de 2º dígito: Eleições Gerais para Presidente no Brasil e a Votação em Urnas Eletrônicas – 2002 a 2018	7
--	---

Cláudia Raquel da Rocha Eirado

Gustavo Pompeu da Silva

Gauss Moutinho Cordeiro

U m Índice Global de Desempenho de Atletas e Equipes de Voleibol	35
---	----

Mateus Duarte de Menezes Ferreira

Thiago Rezende dos Santos

O s Efeitos da Crise Financeira Americana de 2008 na Produção Física Industrial Brasileira: Uma análise de intervenção	51
---	----

Koki Fernando Oikawa

U so de Ferramentas Econométricas para Modelar e Estimar o PIB do Brasil	81
---	----

Waslley Pablo Costa da Silva

Fernando Ferraz do Nascimento

Valmária Rocha da Silva Ferraz

Nota do Editor

É com grande satisfação que anunciamos a publicação do volume 78 (número 243) da Revista Brasileira de Estatística, relativo ao período de Janeiro-Junho de 2020. Este volume traz quatro artigos com variada contribuição, em que são abordados temas envolvendo propostas de metodologia e aplicações diversas.

No primeiro artigo deste número, Cláudia, Gustavo e Gauss exploram a Lei de Benford, que prediz a ocorrência de dígitos, na análise de eleições gerais no Brasil. Os autores utilizam o segundo dígito para avaliar relatórios de eleições de diferentes anos e avaliam a aderência dos resultados ao esperado pela Lei. No segundo artigo Mateus e Thiago trazem um tema de crescente interesse que é o de estatísticas nos esportes. Os autores introduzem um índice de desempenho para o vôlei baseado em fundamentos do esporte. O índice é utilizado para avaliar atletas e equipes em uma temporada. O terceiro artigo de Koki Fernando contextualiza a crise financeira americana de 2008 e seus efeitos e se dedica a analisar, sob a ótica de modelos de séries temporais com intervenção, efeitos na economia brasileira. Fechando este número, Wassley, Fernando e Valmária descrevem e avaliam a performance de diferentes modelos econométricos de séries temporais na estimação do PIB brasileiro de 2019.

É sempre uma satisfação trazer a público mais um número de uma revista de longa história como a RBE, com o irrestrito apoio do IBGE e Associação Brasileira de Estatística, ABE. Me junto aos colegas Editores Marcel de Toledo Vieira e José André de Moura Brito, em agradecimento a autores e revisores, que, anonimamente, contribuíram para mais este número da Revista Brasileira de Estatística.

Desejo a todos uma excelente leitura.

Paulo Justiniano Ribeiro Junior.

Editor Executivo.

ANÁLISE DA LEI DE BENFORD DE 2º DÍGITO: ELEIÇÕES GERAIS PARA PRESIDENTE NO BRASIL E A VOTAÇÃO EM URNAS ELETRÔNICAS – 2002 A 2018

Cláudia Raquel da Rocha Eirado

claudiaeirado@gmail.com

Tribunal Regional Eleitoral do Distrito Federal

Gustavo Pompeu da Silva

gustavopompeu_@hotmail.com

Universidade de São Paulo

Gauss Moutinho Cordeiro

gauus@de.ufpe.br

Universidade Federal de Pernambuco

Resumo: Existe muita especulação em torno da legitimidade dos resultados das Eleições Gerais no Brasil com o uso das urnas eletrônicas. Há *fake news* circulando por toda a Internet em que se levantam suspeitas sobre as urnas e sobre a totalização dos resultados. Para avaliar a confiabilidade das eleições, observaram-se dados públicos das eleições presidenciais brasileiras de 2002 a 2018, obtidos no Repositório de Dados Eleitorais, disponibilizados pelo Tribunal Superior Eleitoral. A distribuição da Lei de Benford é utilizada para detecção de anomalias e fraudes em dados numéricos, observando a frequência com que aparecem os algarismos 0 a 9 nos 1º, 2º, 3º, 4º dígitos das observações. Para dados Eleitorais, recomenda-se observar o padrão dos segundos dígitos da contagem dos votos por município. Utilizou-se o teste Qui-quadrado de Pearson e o teste G (razão de verossimilhança) para avaliar se as distribuições dos segundos dígitos da contagem de votos, para cada candidato, por município, seguem a Lei de Benford – 2BL.

Palavras-Chave: *Eleições; Lei de Benford para o Segundo Dígito – 2BL; Teste de Aderência Qui-quadrado; Teste G (razão de verossimilhança)*

Abstract: There is a lot of speculation around the legitimacy of the results of the General Elections in Brazil with the use of electronic ballot boxes. There are fake news circulating all over the internet where suspicions are raised about the ballot boxes and the results. To evaluate the elections reliability, the public data of the Brazilian Presidential elections from 2002 to 2018 was observed, obtained in the Electoral Data Repository, made available by the *Tribunal Superior Eleitoral*. The Benford's Law distribution is utilized for anomaly and fraud detection in numerical data, observing the frequency in which the digits 0 to 9 appear in the 1st, 2nd, 3rd, 4th digits of the observations. For electoral data, it is recommended to observe the pattern of the 2nd digits of the vote counts by city. It was used the Pearson's Chi-Square test and G-test (likelihood-ratio) to

evaluate if the distributions of the 2nd digits of the vote count for each candidate, by city, follow Benford's Law – 2BL.

Key-Words: *Elections; second-digit Benford's Law; Pearson's Chi-Squared Test; G-test (likelihood-ratio)*

1. INTRODUÇÃO

As urnas eletrônicas brasileiras foram adotadas pela primeira vez em 1996 [2] e foram utilizadas em todo o Estado do Rio de Janeiro, nas capitais dos demais estados e nos municípios com mais de 200 mil eleitores, totalizando 57 cidades no país. Nas eleições de 1998, dois terços do eleitorado votaram eletronicamente. Por fim, no ano 2000, o projeto foi implementado nas Eleições Municipais em sua totalidade no Brasil. Logo, para as Eleições Gerais, as urnas eletrônicas foram utilizadas em todo o território nacional a partir de 2002.

Então, para avaliar a confiabilidade dos resultados das eleições presidenciais no Brasil, observaram-se dados para votação presidencial das Eleições Gerais de 1º e 2º Turnos de 2002, 2006, 2010, 2014 e 2018, disponíveis no site:

<http://www.tse.jus.br/eleicoes/estatisticas/repositorio-de-dados-eleitorais-1>.

A avaliação foi realizada por meio da análise de aderência à Lei de Benford do 2º dígito – 2BL, com o teste de aderência Qui-quadrado de Pearson e o Teste-G (razão de verossimilhança).

A Lei de Newcomb-Benford (LNB) [3] ou Lei de Benford afirma que em dados numéricos que ocorrem naturalmente, o primeiro dígito significativo do número tem maior chance de ser o 1, seguidos de 2, 3, 4, 5, 6, 7, 8 e 9. Por exemplo, em conjuntos que obedecem à lei, o algarismo 1 aparece como o dígito mais significativo em cerca de 30% do tempo, enquanto o 9 aparece como o dígito mais significativo em pouco mais de 4,5% do tempo. Se os dígitos fossem distribuídos uniformemente, cada um deles ocorreria em 11,1% do tempo. A lei de Benford também faz previsões sobre a distribuição de primeiros dígitos – 1BL, segundos dígitos – 2BL, terceiros dígitos – 3BL, combinações de dígitos e assim por diante.

Segue então que, a probabilidade de um algarismo d ser encontrado da n ésima posição (n) de um número é dada por:

$$p(d) = \sum_{k=10^{n-2}}^{10^{n-1}-1} \log_{10} \left(1 + \frac{1}{10k+d} \right), \quad (1)$$

em que $d = \{0, 1, 2, 3, 4, 5, 6, 7, 8, 9\}$.

Por exemplo, a probabilidade do algarismo 3 ser o segundo dígito de um número qualquer é igual a:

$$\log_{10} \left(1 + \frac{1}{13} \right) + \log_{10} \left(1 + \frac{1}{23} \right) + \dots + \log_{10} \left(1 + \frac{1}{93} \right) \approx 0,10433$$

A lei então segue as seguintes probabilidades:

Tabela 1 - Frequências esperadas base na Lei de NewcombBenford

Algarismo	1ª posição	2ª posição	3ª posição	4ª posição
0	-	0,11968	0,10178	0,10018
1	0,30103	0,11389	0,10138	0,10014
2	0,17609	0,10882	0,10097	0,10010
3	0,12494	0,10433	0,10057	0,10006
4	0,09691	0,10031	0,10018	0,10002
5	0,07918	0,09668	0,09979	0,09998
6	0,06695	0,09337	0,09940	0,09994
7	0,05799	0,09035	0,09902	0,09990
8	0,05115	0,08757	0,09864	0,09986
9	0,04576	0,08500	0,09827	0,09982

Por exemplo, para a sequência $S = \{2^n \mid n=4, \dots, 14\} = \{16, 32, 64, 128, 256, 512, 1024, 2048, 4096, 8192, 16384\}$, a sequência dos primeiros dígitos será 1BL = {1, 3, 6, 1, 2, 5, 1, 2, 4, 8, 1}. A sequência dos segundos dígitos será 2BL = {6, 2, 4, 2, 5, 1, 0, 0, 0, 1, 6}.

A lei tem sido aplicada em dados como contas de eletricidade, endereços, preços de ações, preços de casas, números populacionais, taxas de mortalidade, comprimentos de rios, constantes físicas e matemáticas, e processos descritos com leis de potência – que são muito comuns na natureza. Ela tende a ser mais precisa quando os valores são distribuídos em várias ordens de magnitude. Entretanto, não funciona para lista de telefones, por exemplo, pois os números iniciais costumam ser determinados e fixados pelas operadoras, então estas sequências não ocorrem de forma espontânea.

Existem discussões sobre a aplicabilidade da LNB para dados eleitorais. Mebane (2006) [1] afirma que uma forma de analisar se há fraude nos resultados de eleições é verificar o comportamento da distribuição dos segundos dígitos das contagens dos votos por cidades ou municípios. A análise por seção de votação, que só possui uma urna, não é eficaz, pois o universo amostral é muito pequeno e a Lei de Benford é uma distribuição limite (assintótica). No caso das seções eleitorais brasileiras pode haver entre 1 e 500 eleitores e a contagem de votos não possui várias ordens de magnitude. O mesmo ocorre por zonas eleitorais, pois a magnitude é limitada. Tanto no caso das seções eleitorais como das zonas eleitorais é a Justiça Eleitoral que define a quantidade nesses agrupamentos,

conforme determinam a Resolução TSE 23.422, de 6 de maio de 2014 e a Resolução TSE 23.520, de 1º de junho de 2017. Por outro lado, a análise por municípios possui muitas observações com diversas ordens de magnitude e a quantidade de eleitores surge de forma espontânea.

Segundo Brady (2005) [4], sobre as eleições americanas, certamente pode-se imaginar que os processos estocásticos que resultam em uma corrida competitiva de dois candidatos em distritos com magnitudes que variam entre 100 e 1000 eleitores, o primeiro dígito modal para cada voto do candidato não será 1 ou 2, mas sim 4, 5 ou 6, pois cada um dos dois candidatos receberá aproximadamente entre 400 e 600 votos. É precisamente este exemplo que o leva a rejeitar a 1BL como um teste de fraude e um motivo para se concentrar, se necessário, em 2BL. Por analogia, funcionaria da mesma forma para os dados brasileiros.

O que é preciso saber é se a distribuição dos dígitos da contagem da votação eletrônica segue a Lei de Benford do segundo dígito – 2BL.

Estudos apontam que dados populacionais municipais seguem a distribuição de LNB [5]. Então, considerando que o número de eleitores deriva da população municipal e que o número de votos está altamente correlacionado com o número de eleitores, pode-se imaginar que os dígitos da contagem da votação eletrônica podem seguir uma distribuição desse tipo.

2. METODOLOGIA

Para verificar a aderência da distribuição dos dados com a Lei de Benford – 2BL, pode-se utilizar o teste Qui-quadrado de Pearson. O teste avalia a qualidade do ajuste e estabelece se uma distribuição de frequências observadas (nesse estudo, serão as frequências dos dígitos das contagens de votos) difere de uma distribuição teórica ou das frequências esperadas (no caso concreto, serão as probabilidades da Lei de Benford). A estatística do teste é dada por:

$$Q = \sum_{i=1}^k \frac{(O_i - E_i)^2}{E_i}, \quad (2)$$

Em que

O_i = frequências observadas na categoria i .

E_i = frequências esperadas na categoria i .

k é o número de categorias observadas.

A estatística possui, em grandes amostras, a distribuição Qui-quadrado com $(k - 1)$ graus de liberdade, isto é, $Q \sim \chi^2_{k-1}$. Nesse teste, não rejeitar a hipótese nula significa que a distribuição observada e a distribuição esperada não diferem. Ao rejeitar a hipótese nula, aceita-se a hipótese alternativa de que as distribuições observadas e esperadas (modelo teórico – 2BL) são diferentes. No caso desse estudo, tem-se $k = 10$, pois são 10 dígitos (0 a 9), logo $Q \sim \chi^2_9$.

O Teste-G, recomendado por Sokal & Rohlf (1981) [8] é alternativa ao Teste de Aderência Qui-quadrado de Pearson e é mais eficiente que o primeiro pelo critério de Badahur, e são igualmente eficientes no sentido de Pitman ou no sentido de Hodges e Lehmann, de acordo com Quine & Robinson (1985) [9]. Os Testes-G são testes de significância estatística de razão de verossimilhança.

Testes de razão de verossimilhança (RV) são capazes de verificar a bondade do ajuste de dois modelos estatísticos pela razão de suas probabilidades. Intrinsecamente compara-se o modelo irrestrito (com a maximização em todo o espaço paramétrico Θ) e aquele restrito com subconjunto Θ_0 do espaço paramétrico Θ , isto é, $\Theta_0 \in \Theta$. A hipótese nula é $H_0 = \theta \in \Theta_0$ e a hipótese alternativa é $H_1 = \theta \in \Theta \setminus \Theta_0$ (*complementar de Θ_0*). A estatística do teste sob H_0 é dada por:

$$\lambda_{RV} = -2 \ln \left[\frac{\sup_{\theta \in \Theta_0} \mathcal{L}(\theta)}{\sup_{\theta \in \Theta} \mathcal{L}(\theta)} \right], \quad (3)$$

Em que $\mathcal{L}(\theta)$ é a função de máxima verossimilhança. Assintoticamente, $\lambda_{RV} \sim \chi^2_b$, em que b (graus de liberdade) é a diferença de dimensionalidade entre Θ e Θ_0 .

Os testes Qui-quadrado são aproximações da razão de log-verossimilhança dos testes G. No estudo em questão, o teste G compara quão próximas estão as probabilidades esperadas pela Lei de Benford do 2º dígito e aquelas observadas pelos dados eleitorais. A fórmula geral do teste é:

$$G = 2 \sum_i O_i \ln \left(\frac{O_i}{E_i} \right), \quad (4)$$

em que $\sum_i O_i = \sum_i E_i = N$. A estatística G tem distribuição assintótica Qui-quadrado com os mesmos $(k - 1)$ graus de liberdade do teste Qui-quadrado, ou seja, $G \sim \chi^2_{k-1}$, então $G \sim \chi^2_9$.

2.1 ANÁLISE DE DADOS ELEITORAIS – 2002 A 2018

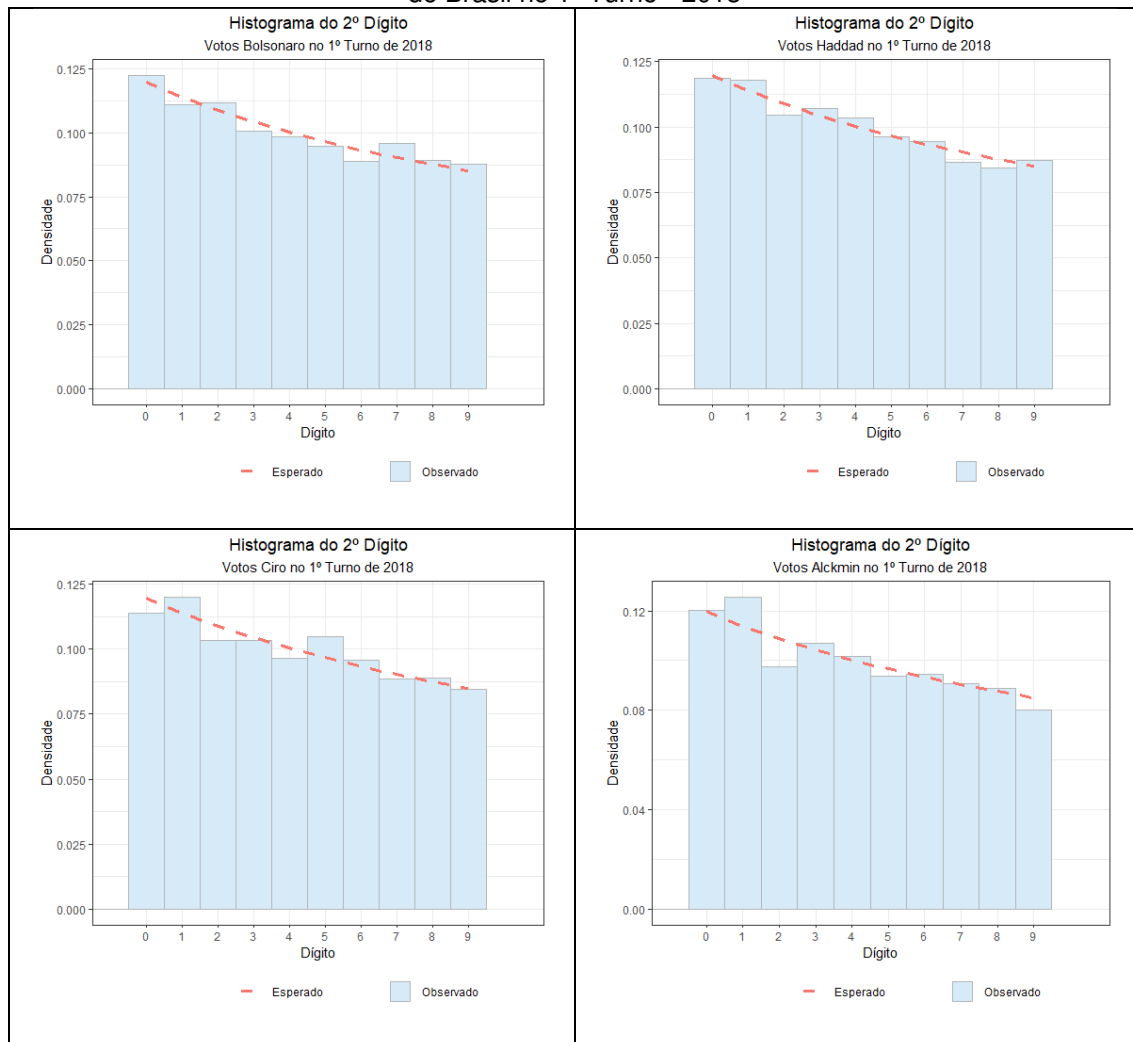
Partiu-se dos dados mais recentes das Eleições de 2018 disponíveis em: http://agencia.tse.jus.br/estatistica/sead/odsele/votacao_secao/votacao_secao_2018_BR.zip.

Tomaram-se as contagens de votação por candidatos agregadas por municípios do Brasil e observaram-se os 2º dígitos dessas contagens.

No 1º Turno foram analisados os cinco primeiros colocados: Bolsonaro, Haddad, Ciro, Alckmin e Amoêdo. Para o 2º Turno: Bolsonaro x Haddad.

A seguir, apresenta-se o resultado para o teste da **Lei de Benford do 2º Dígito (2BL)**, em 2018:

Figura 1 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 1º Turno - 2018



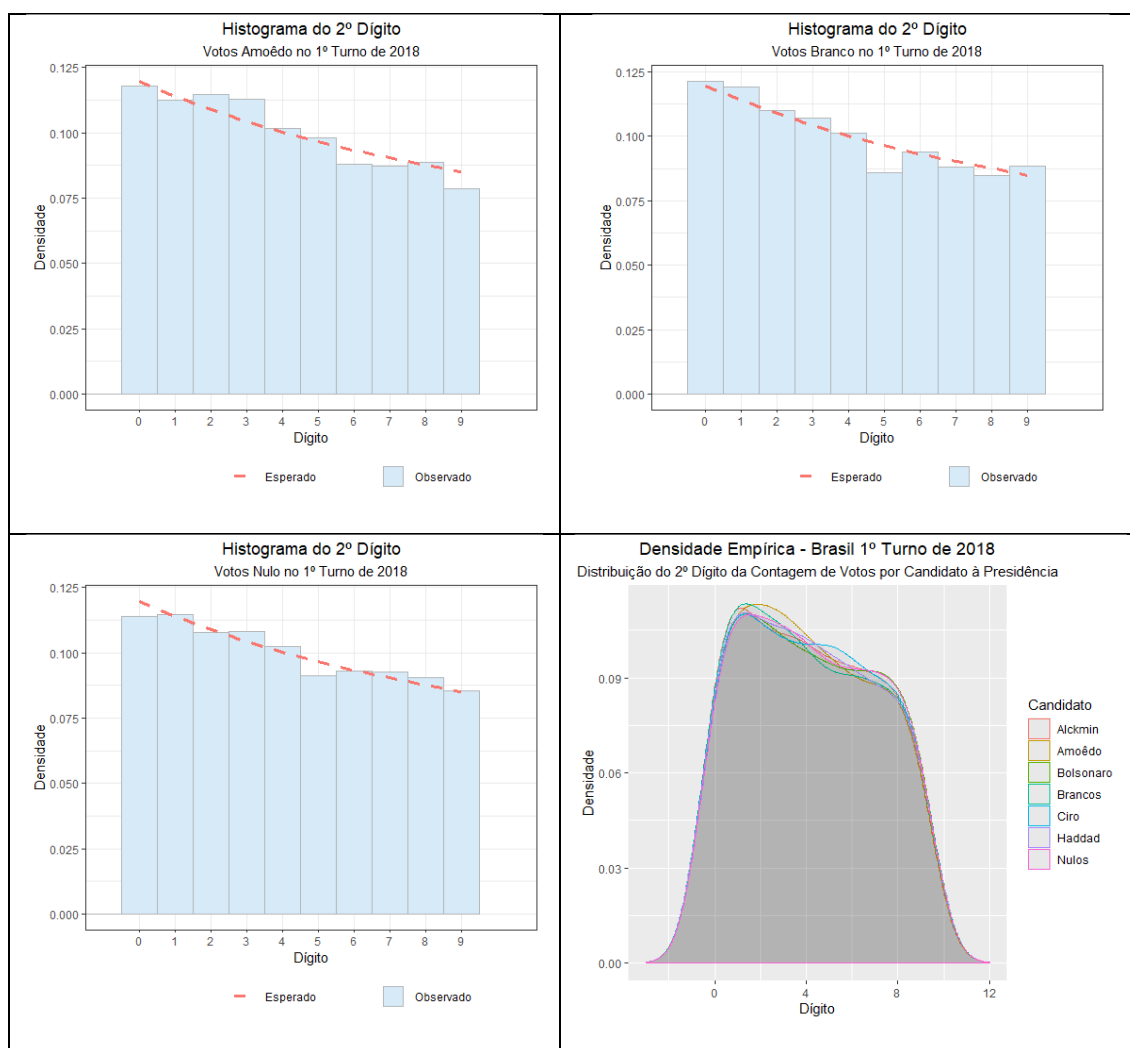


Tabela 2 - Teste Qui-Quadrado - 2018 - 1º Turno - Teste 2BL - Dígitos 0 a 9

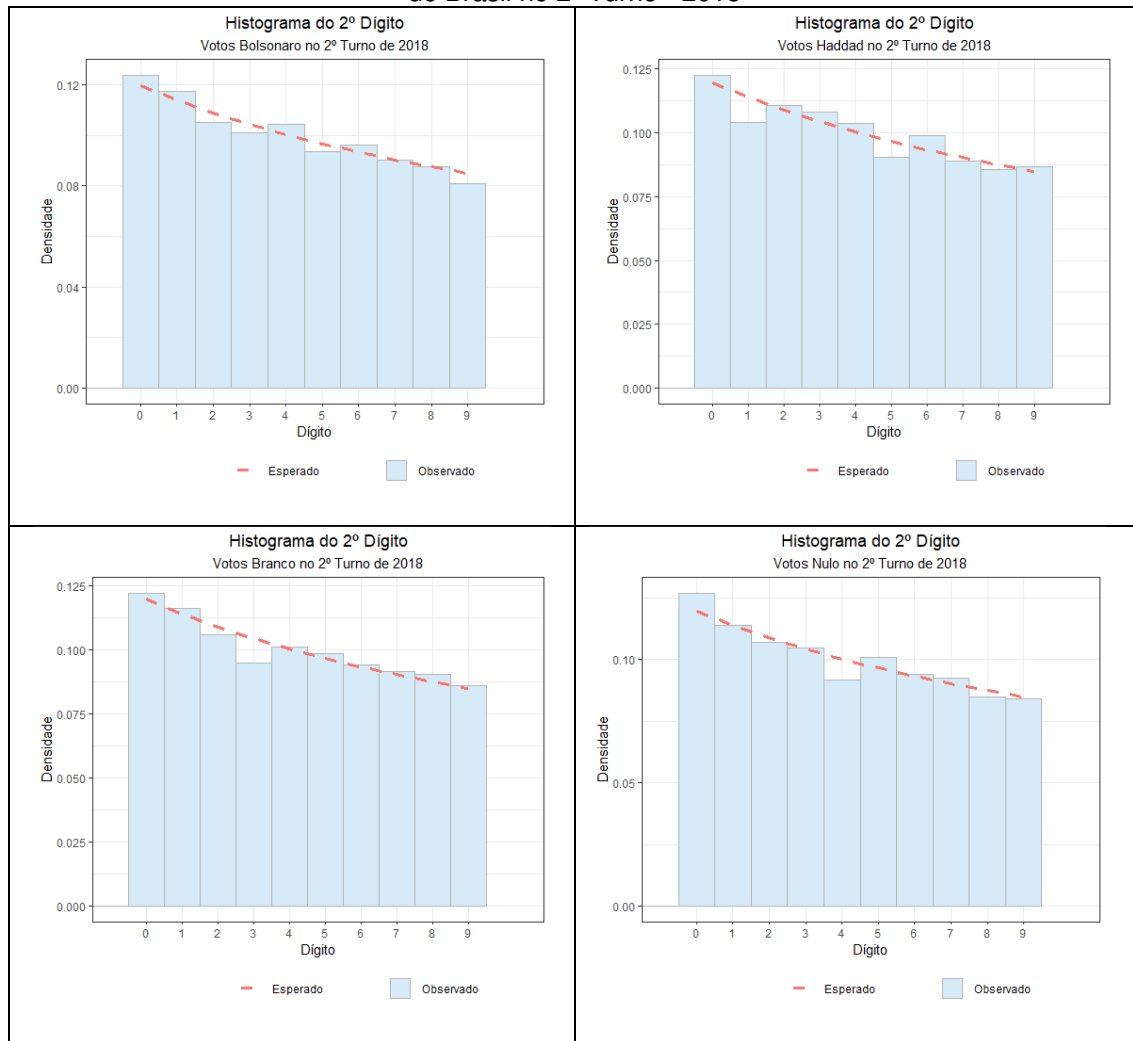
Candidato	Estatística χ^2	gl	p-valor
Bolsonaro	6,0616	9	0,7337
Haddad	4,7813	9	0,8529
Ciro	10,26	9	0,3299
Alckmin	15,995	9	0,06699
Amoedo	9,7767	9	0,3689
Brancos	10,166	9	0,3372
Nulos	5,3199	9	0,8056

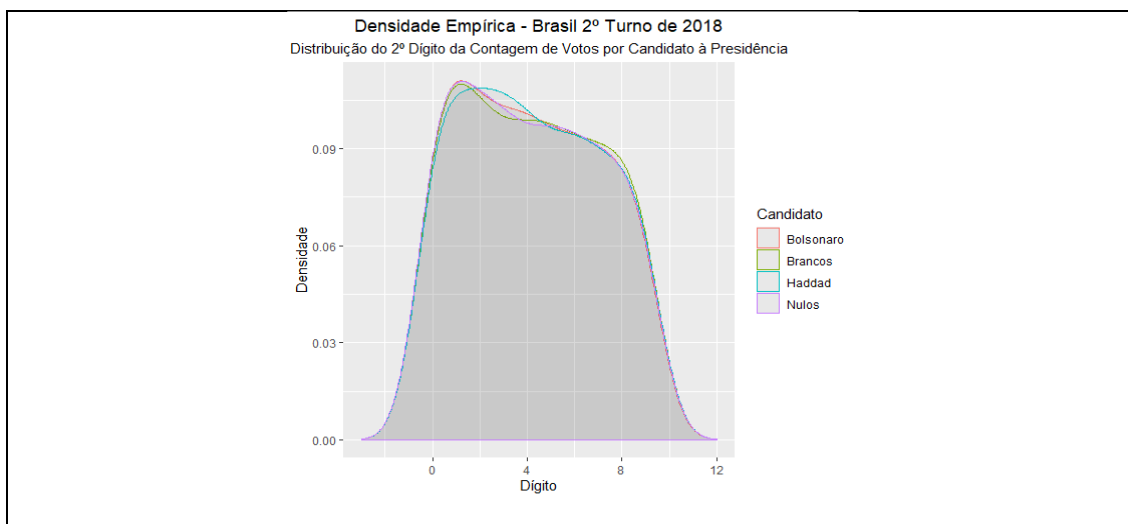
Tabela 3 - Teste G (Razão de Verossimilhança) - 2018 - 1º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Bolsonaro	6,0522	9	0,7347
Haddad	4,7981	9	0,8515
Ciro	10,181	9	0,3361
Alckmin	16,048	9	0,06589
Amoedo	9,7552	9	0,3707
Brancos	10,394	9	0,3196
Nulos	5,3553	9	0,8023

No 1º Turno de 2018, o teste Qui-Quadrado e o teste G para 2BL, com nível de significância de 5% ou nível de confiança de 95%, apontam que não existe diferença significativa entre a distribuição observada para os candidatos Bolsonaro, Haddad, Ciro, Alckmin, Amoedo, votos brancos e votos nulos e aquela da Lei de Benford para o 2º dígito. Portanto, não existem evidências de fraudes nas urnas.

Figura 2 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 2º Turno - 2018



**Tabela 4** - Teste Qui-Quadrado - 2018 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística χ^2	gl	p-valor
Bolsonaro	5,8553	9	0,7543
Haddad	11,542	9	0,2404
Brancos	6,7977	9	0,6582
Nulos	8,6626	9	0,4690

Tabela 5 - Teste G (Razão de Verossimilhança) - 2018 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Bolsonaro	5,8669	9	0,7532
Haddad	11,675	9	0,2323
Brancos	6,9468	9	0,6427
Nulos	8,7265	9	0,4629

No 2º Turno de 2018, o teste Qui-Quadrado e o teste G (Razão de Verossimilhança) para 2BL, com nível de significância de 5%, indicam que não existe diferença significativa entre a distribuição observada para Bolsonaro, Haddad, votos brancos e votos nulos e aquela dada pela Lei de Benford para o 2º dígito. Portanto, não há indícios de fraude na contagem de votos.

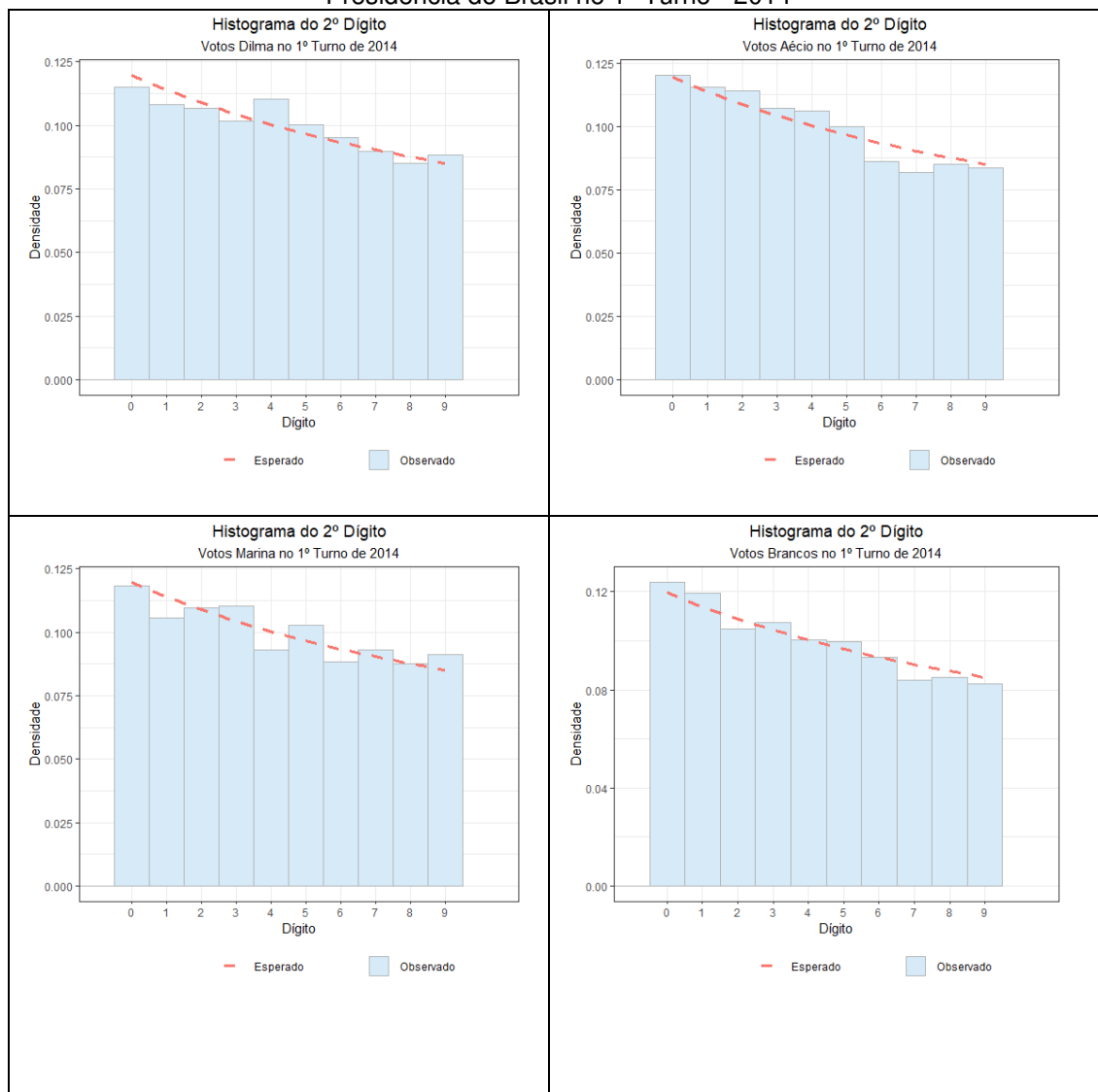
Os dados de votação para presidente em 2014 foram retirados do [link: http://agencia.tse.jus.br/estatistica/sead/odsele/votacao_secao/votacao_secao_2014_BR.zip](http://agencia.tse.jus.br/estatistica/sead/odsele/votacao_secao/votacao_secao_2014_BR.zip).

Foram realizados os testes da Lei de Newcomb-Benford (LNB) para o 2º dígito (2BL). Para verificar se a distribuição dos dados segue a LNB, utilizou-se o teste de aderência Qui-quadrado.

No 1º Turno foram analisados os três primeiros colocados: Aécio, Dilma e Marina. Para o 2º Turno: Dilma x Aécio.

A seguir, apresenta-se o resultado para o teste da **Lei de Benford do 2º Dígito (2BL)**, em 2014:

Figura 3 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 1º Turno - 2014



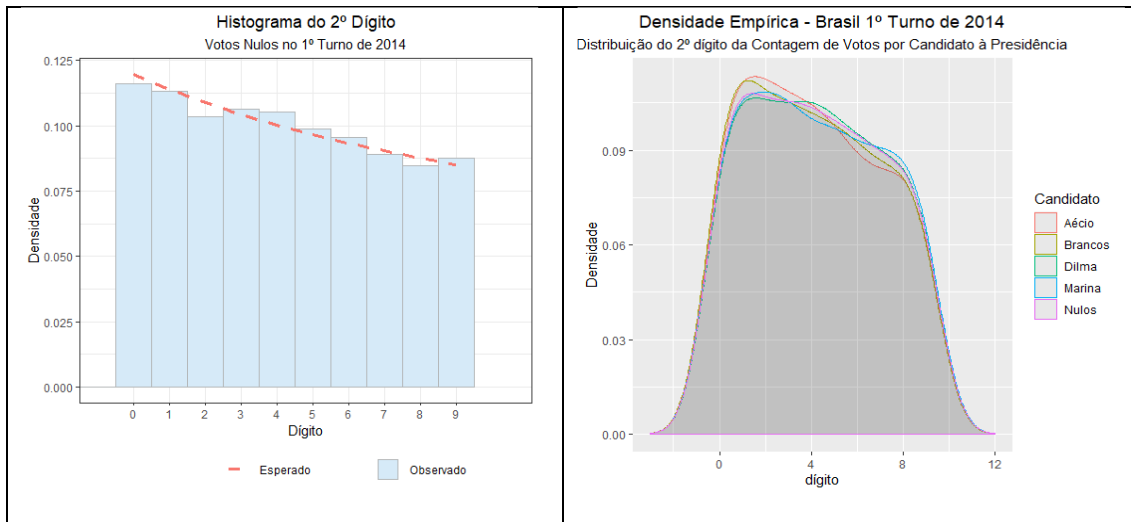


Tabela 6 - Teste Qui-Quadrado - 2014 - 1º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística χ^2	gl	p-valor
Dilma	11,412	9	0,2485
Aécio	12,779	9	0,1729
Marina	15,224	9	0,08496
Brancos	7,5639	9	0,5786
Nulos	5,2592	9	0,8112

Tabela 7 - Teste G (Razão de Verossimilhança) - 2014 - 1º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Dilma	11,26	9	0,2583
Aécio	12,942	9	0,1652
Marina	15,27	9	0,08378
Brancos	7,6038	9	0,5745
Nulos	5,2634	9	0,8108

No 1º Turno de 2014, tanto o teste Qui-Quadrado quanto o teste G (Razão de Verossimilhança) para 2BL, com nível de significância de 5%, indicam que não existe diferença significativa entre a distribuição observada para Aécio, Dilma, Marina, votos brancos e votos nulos e aquela preconizada pela Lei de Benford para o 2º dígito. Não existem indícios de fraude na contagem de votos.

Figura 4 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 2º Turno - 2014

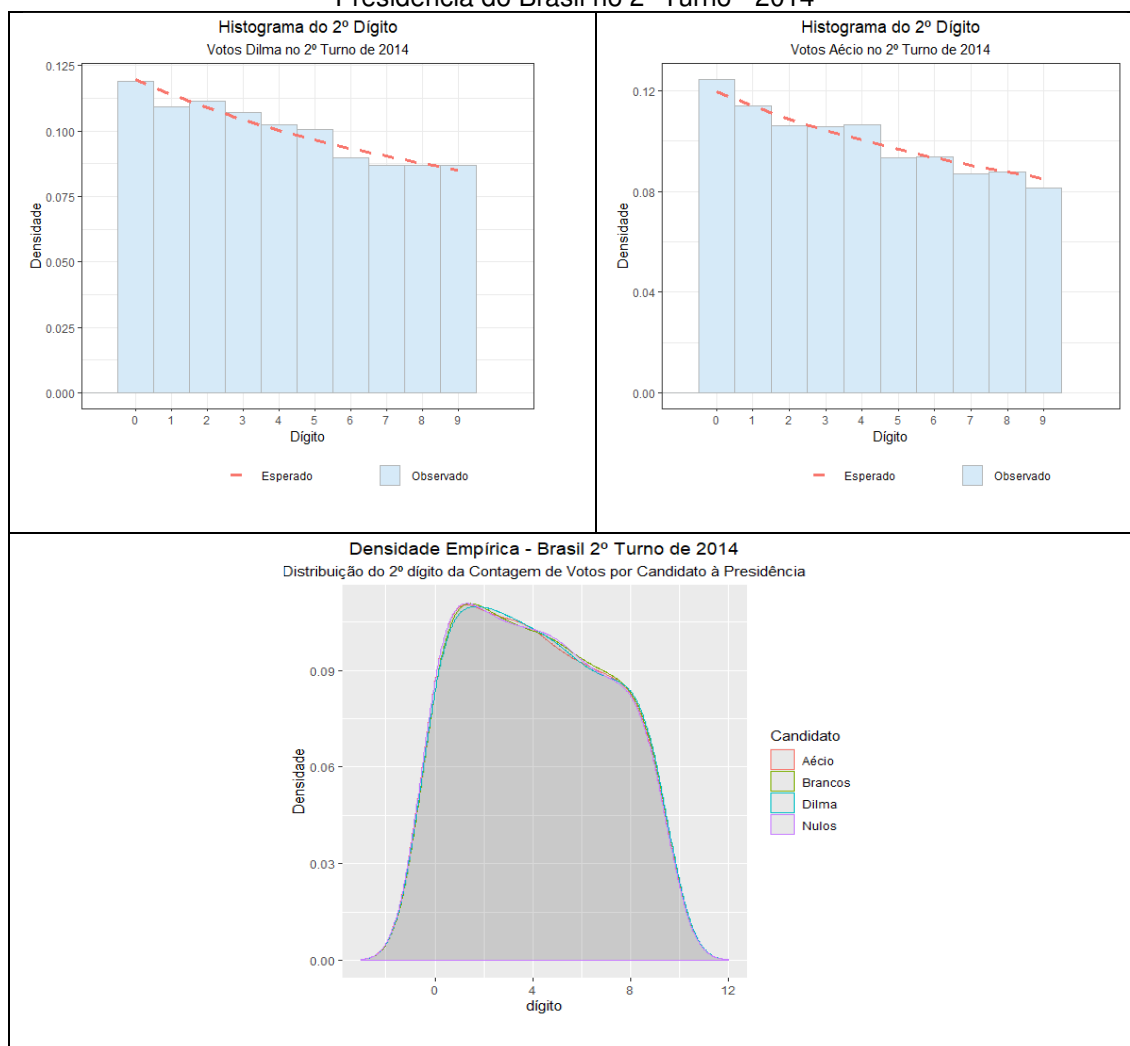


Tabela 8 - Teste Qui-Quadrado - 2014 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística χ^2	gl	p-valor
Dilma	4,784	9	0,8527
Aécio	6,0361	9	0,7363
Brancos	4,2592	9	0,8935
Nulos	6,44	9	0,6952

Tabela 9 - Teste G (Razão de Verossimilhança) - 2014 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Dilma	4,7983	9	0,8515
Aécio	6,0085	9	0,7391
Brancos	4,2575	9	0,8937
Nulos	6,4255	9	0,6967

No 2º Turno de 2014, o teste Qui-Quadrado e o teste G (Razão de Verossimilhança) para 2BL, com nível de significância de 5%, indica que não existe diferença significativa entre a distribuição observada para Aécio, Dilma,

votos brancos e votos nulos e aquela dada pela Lei de Benford para o 2º dígito. E, portanto, não há indícios de fraude.

No Apêndice I, apresentam-se os testes Qui-quadrado e os testes G (Razão de Verossimilhanças) com 2BL para as votações das eleições de 2010, 2006 e 2002, respectivamente. Em todos eles, a conclusão é de que o teste 2BL indica que não existem evidências de fraudes na contagem de votos nas eleições brasileiras no período pós urnas eletrônicas. O programa desenvolvido em **R** [7] é apresentado no Apêndice II.

3. CONCLUSÃO

Para o teste Qui-quadrado e para o Teste G (Razão de Verossimilhança) 2BL, com nível de confiança de 95%, conclui-se que não existem evidências de fraudes na contagem de votos nas Eleições Gerais de 1º e 2º Turnos de 2002, 2006, 2010, 2014 e 2018.

Mebane et al (2018) [6] também indica que não existem evidências de fraudes nas Eleições Gerais de 2014 no Brasil por meio de outra metodologia de Estimção Bayesiana de FMM (Finite Mixture Models) para fraudes eleitorais (incluindo covariáveis demográficas e sócio-econômicas).

4. REFERÊNCIAS BIBLIOGRÁFICAS

- [1] **Mebane**, W. R. Jr. (2006). "Election Forensics: The Second-digit Benford's Law Test and Recent American Presidential Elections". *Prepared for delivery at the Election Fraud Conference, Salt Lake City, Utah, September 29--30, 2006*. Disponível em <http://www-personal.umich.edu/~wmebane/fraud06.pdf>, em 17/10/2018.
- [2] **Brasil, Tribunal Superior Eleitoral**. (2016). Urna eletrônica: 20 anos a favor da democracia. – Brasília: Tribunal Superior Eleitoral, 2016. 44 p.; 24 cm.
- [3] **Benford**, F. (1938). "The law of anomalous numbers". *Proc. Am. Philos. Soc.* 78 (4): 551–572. JSTOR 984802
- [4] **Brady**, H. E. (2005). "Comments on Benford's Law and the Venezuelan Election." Unpublished manuscript, Stanford University, January 19, 2005.
- [5] **Cinelli**, C. (2014) "Indício de fraude nas eleições? Usando a Lei de Benford." Disponível em: <https://analisereal.com/2014/11/02/indicio-de-fraude-nas-eleicoes-usando-a-lei-de-benford/>. Acesso em 17/10/2018.
- [6] **Ferrari, D., McAlister, K. and Mebane, W. R. Jr.** (2018). "Developments in Positive Empirical Models of Election Frauds: Dimensions and Decisions". Prepared for presentation at the 2018 Annual Meeting of the Society for Political Methodology, Provo, UT, July 16--18.
- [7] **R Core Team**. (2016) "R: A Language and Environment for Statistical Computing". R Foundation for Statistical Computing, Vienna, Austria.
- [8] **Sokal, R. R.; Rohlf, F. J.** (1981). *Biometry: The Principles and Practice of Statistics in Biological Research* (Second ed.). New York: Freeman. Chapter 17.
- [9] **Quine, M. P.; Robinson, J.** (1985). "Efficiencies of chi-square and likelihood ratio goodness-of-fit tests". *Annals of Statistics*. 13 (2): 727–742.

APÊNDICE 1

A seguir, apresenta-se o resultado para o teste da **Lei de Benford do 2º Dígito (2BL)**, em 2010:

Figura 5 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 1º Turno - 2010

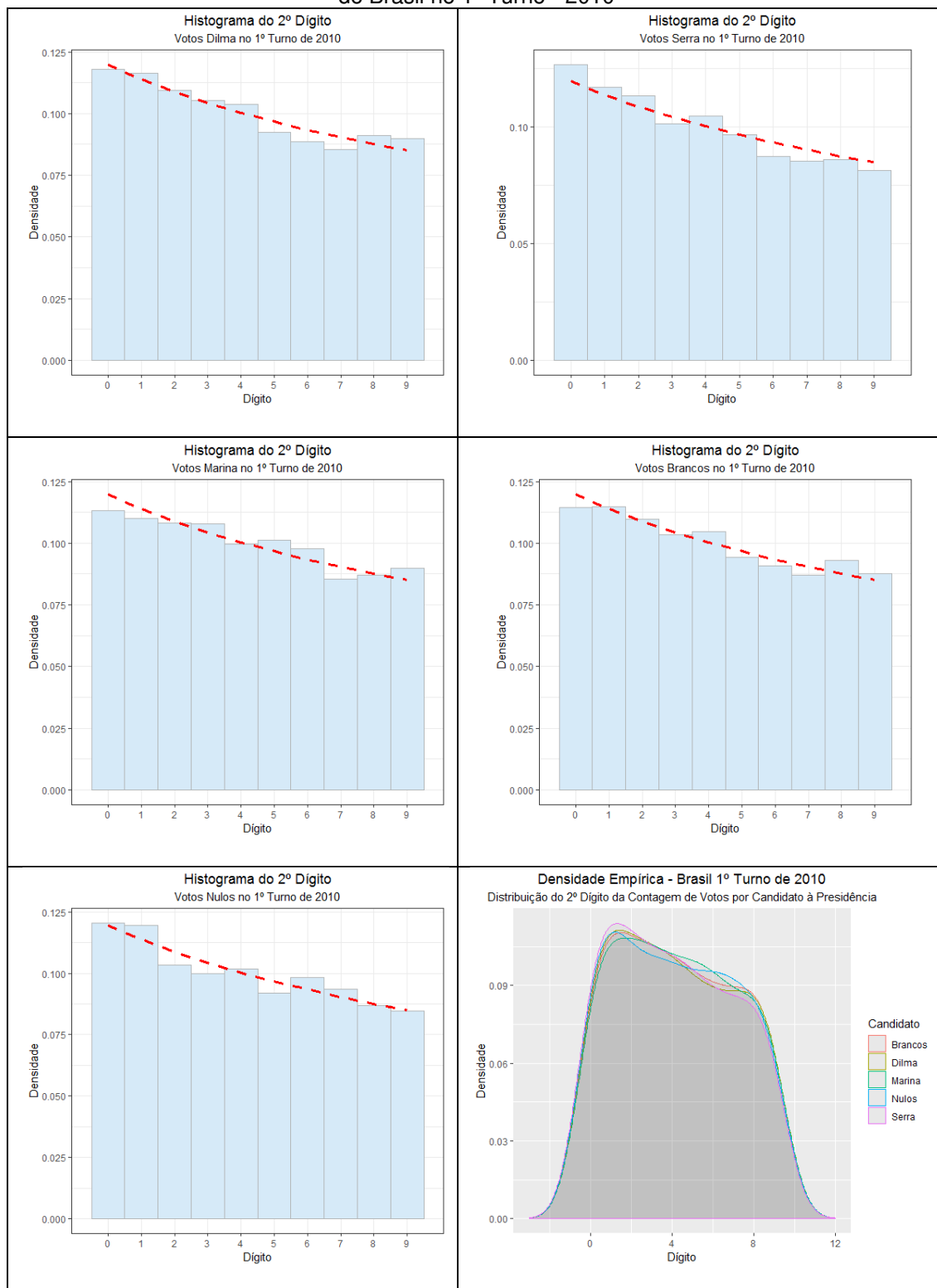
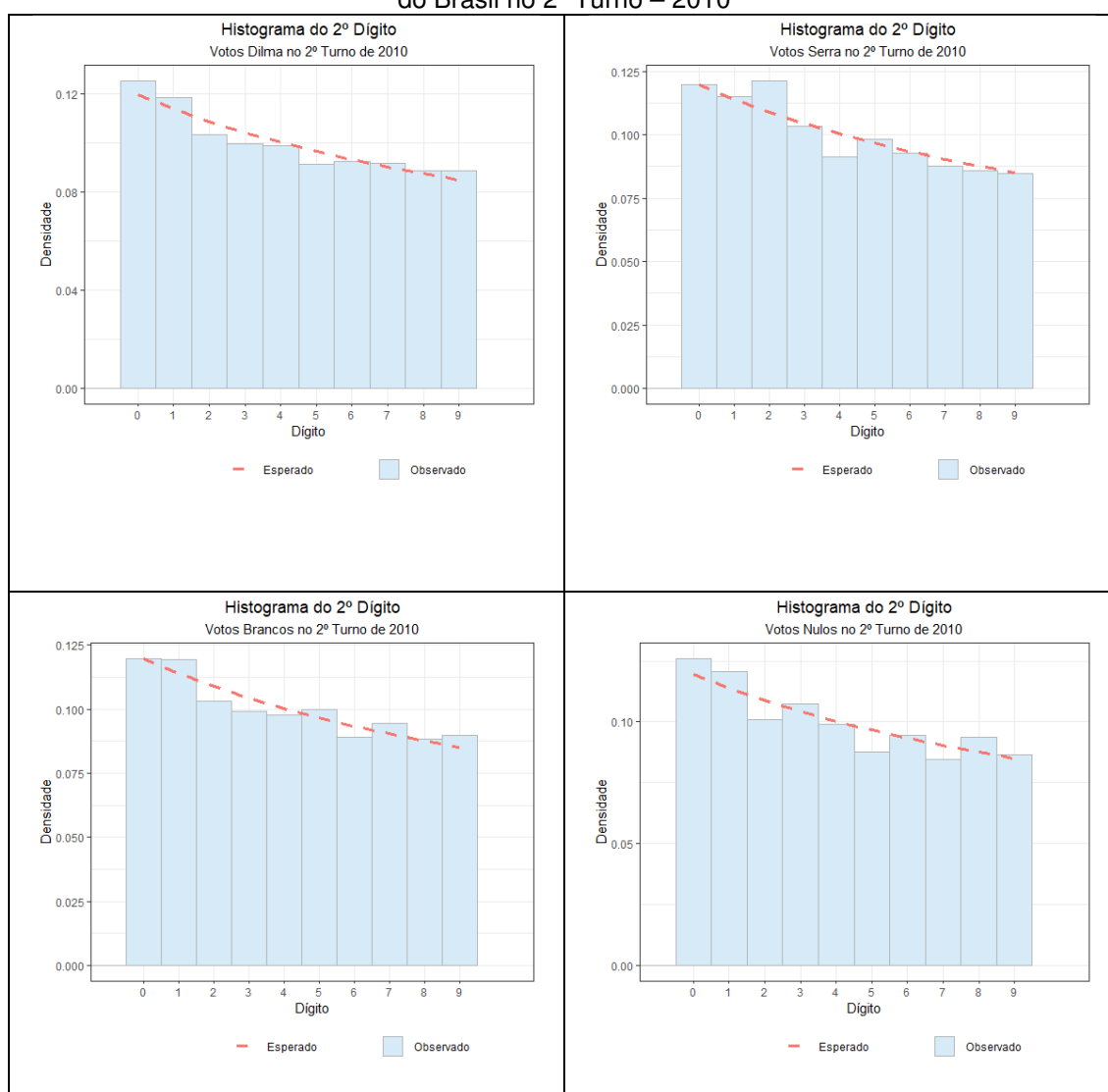


Tabela 10 - Teste Qui-Quadrado - 2010 - 1º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística χ^2	gl	p-valor
Dilma	7,4432	9	0,5911
Serra	9,9592	9	0,3538
Marina	8,9973	9	0,4375
Branco	6,2113	9	0,7186
Nulos	7,7271	9	0,5619

Tabela 11 - Teste G (Razão de Verossimilhança) - 2010 - 1º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Dilma	7,4643	9	0,5889
Serra	9,9758	9	0,3524
Marina	8,9986	9	0,4374
Branco	6,1897	9	0,7208
Nulos	7,733	9	0,5613

Figura 6 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 2º Turno – 2010


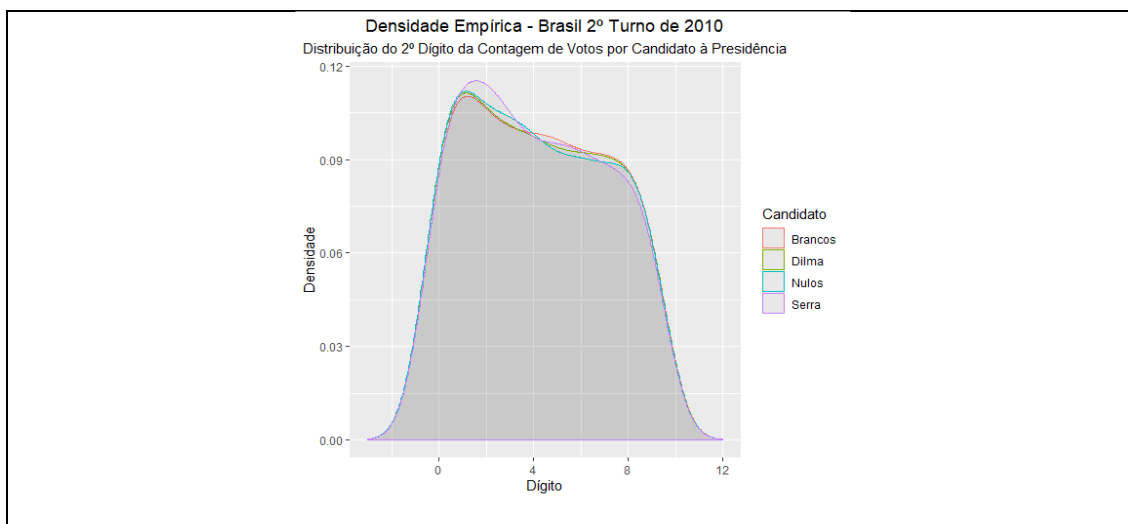


Tabela 12 - Teste Qui-Quadrado - 2010 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística χ^2	gl	p-valor
Dilma	8,2465	9	0,5095
Serra	13,1	9	0,1581
Brancos	9,3337	9	0,4071
Nulos	17,207	9	0,04557

Tabela 13 - Tabela 7.2 - Teste G (Razão de Verossimilhança) - 2010 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Dilma	8,2636	9	0,5078
Serra	12,955	9	0,1647
Brancos	9,3389	9	0,4066
Nulos	17,367	9	0,04327

A seguir, apresenta-se o resultado para o teste da **Lei de Benford do 2º Dígito (2BL)**, em 2006:

Figura 7 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 1º Turno - 2006

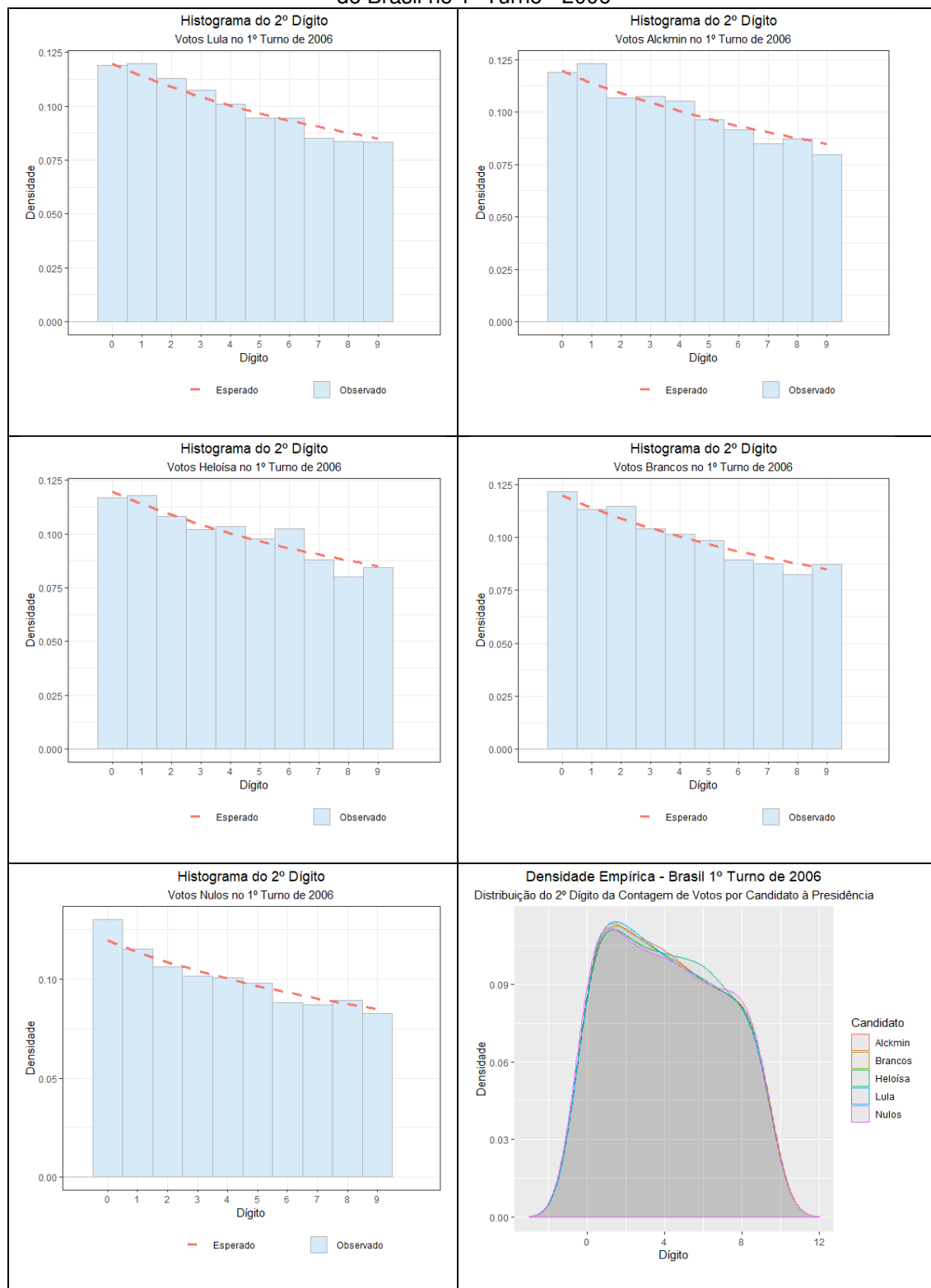


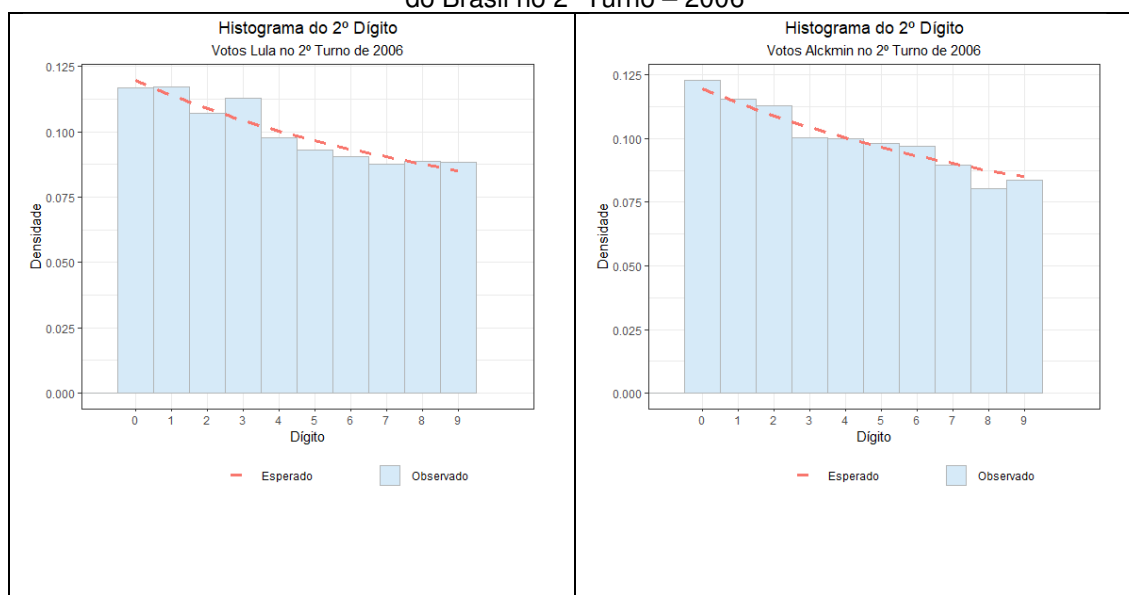
Tabela 14 - Teste Qui-Quadrado - 2006 - 1º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística χ^2	gl	p-valor
Lula	6,038	9	0,7361
Alckmin	10,043	9	0,347
Heloísa	11,04	9	0,273
Branços	5,7162	9	0,768
Nulos	8,8414	9	0,452

Tabela 15 - Teste G (Razão de Verossimilhança) - 2006 - 1º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Lula	6,0481	9	0,7351
Alckmin	9,9991	9	0,3506
Heloísa	11,003	9	0,2755
Branços	5,7369	9	0,7659
Nulos	8,7382	9	0,4618

Figura 8 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 2º Turno – 2006



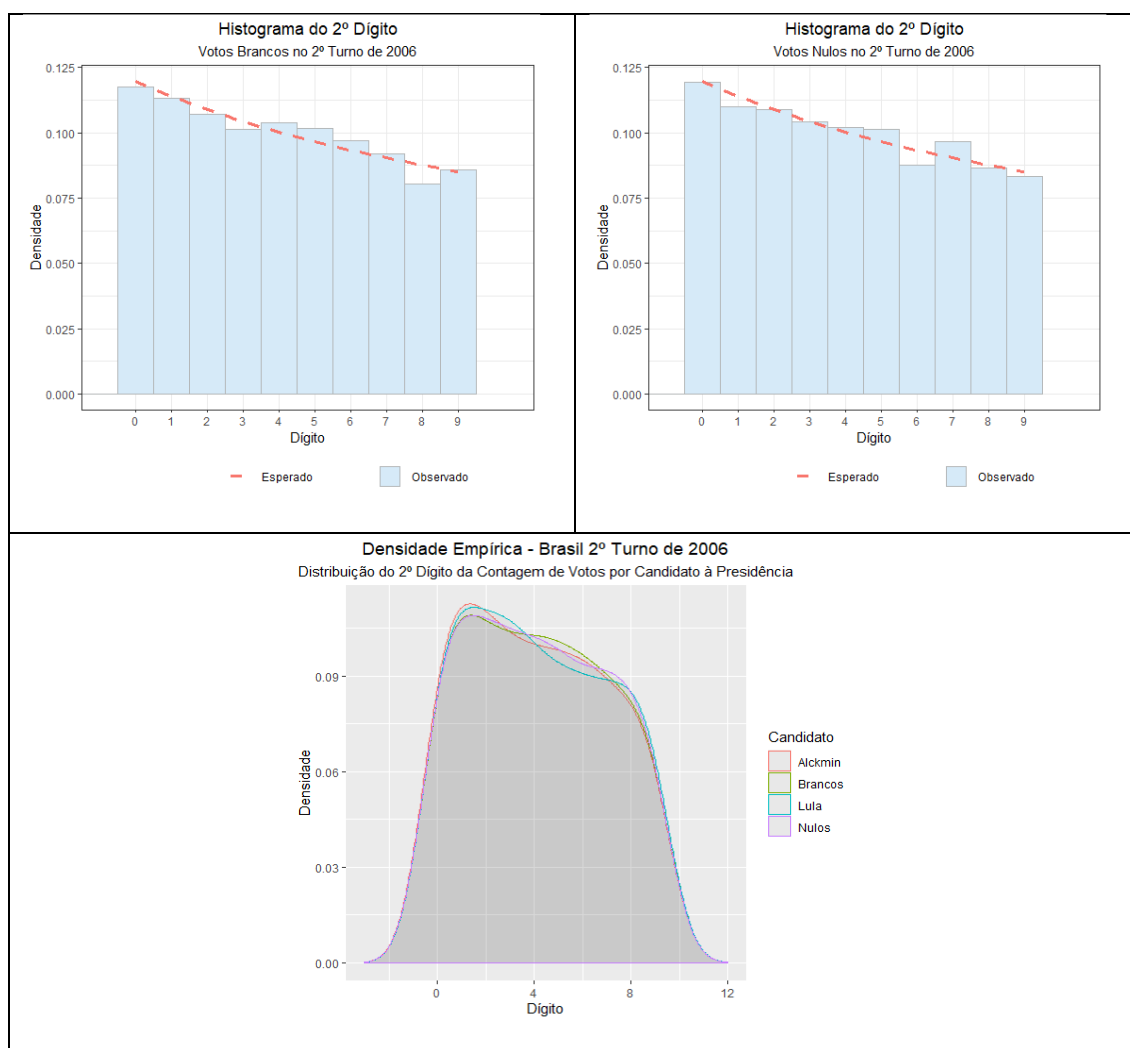


Tabela 16 - Teste Qui-Quadrado - 2006 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística χ^2	gl	p-valor
Lula	7,9464	9	0,5396
Alckmin	6,9441	9	0,6429
Brancos	7,2632	9	0,6097
Nulos	6,6536	9	0,6731

Tabela 17 - Teste G (Razão de Verossimilhança) - 2006 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Lula	7,8523	9	0,5491
Alckmin	7,0297	9	0,634
Brancos	7,3197	9	0,6039
Nulos	6,6303	9	0,6755

A seguir, apresenta-se o resultado para o teste da **Lei de Benford do 2º Dígito (2BL)**, em 2002:

Figura 9 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 1º Turno - 2002

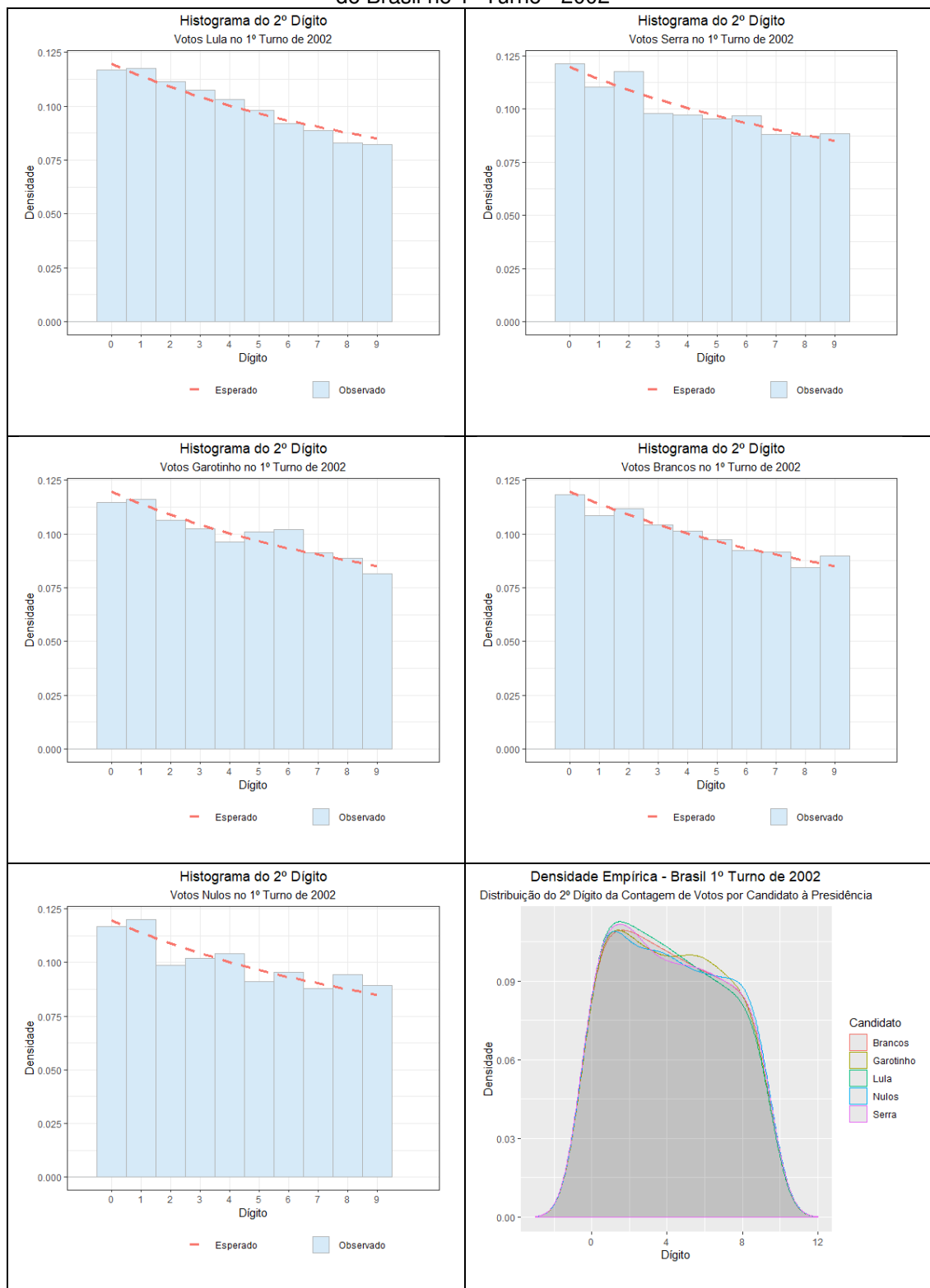


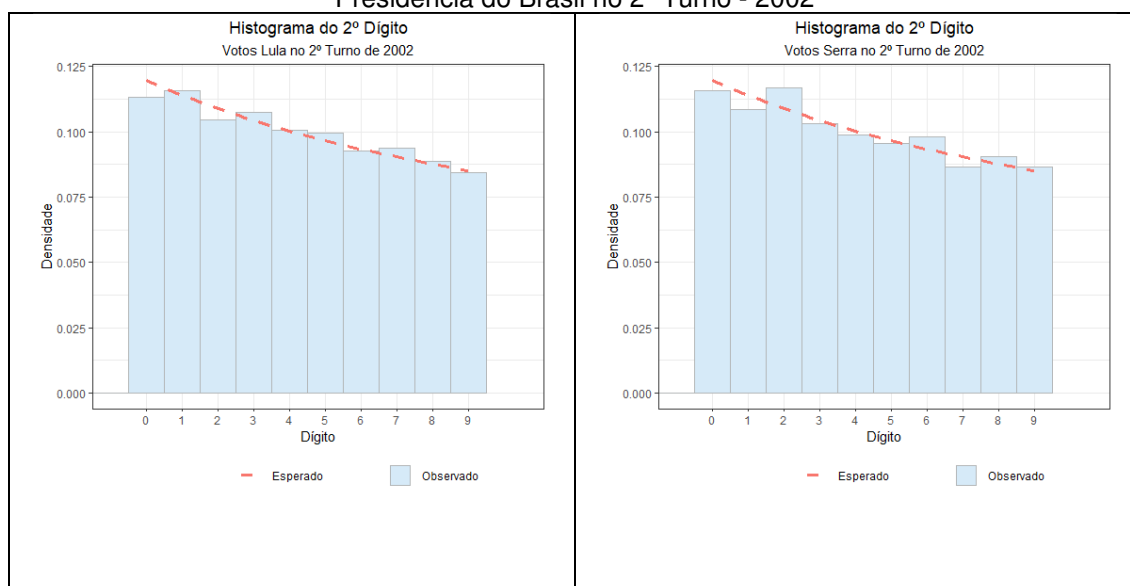
Tabela 18 - Teste Qui-Quadrado - 2002 - 1º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística χ^2	gl	p-valor
Lula	4,6426	9	0,8643
Serra	9,3068	9	0,4094
Garotinho	9,3424	9	0,4063
Branços	4,4347	9	0,8805
Nulos	15,355	9	0,08163

Tabela 19 - Teste G (Razão de Verossimilhança) - 2002 - 1º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Lula	4,6609	9	0,8628
Serra	9,2551	9	0,4141
Garotinho	9,235	9	0,4159
Branços	4,4315	9	0,8808
Nulos	15,418	9	0,08007

Figura 10 – Distribuições dos 2º Dígitos das Contagens de Votos dos Candidatos à Presidência do Brasil no 2º Turno - 2002



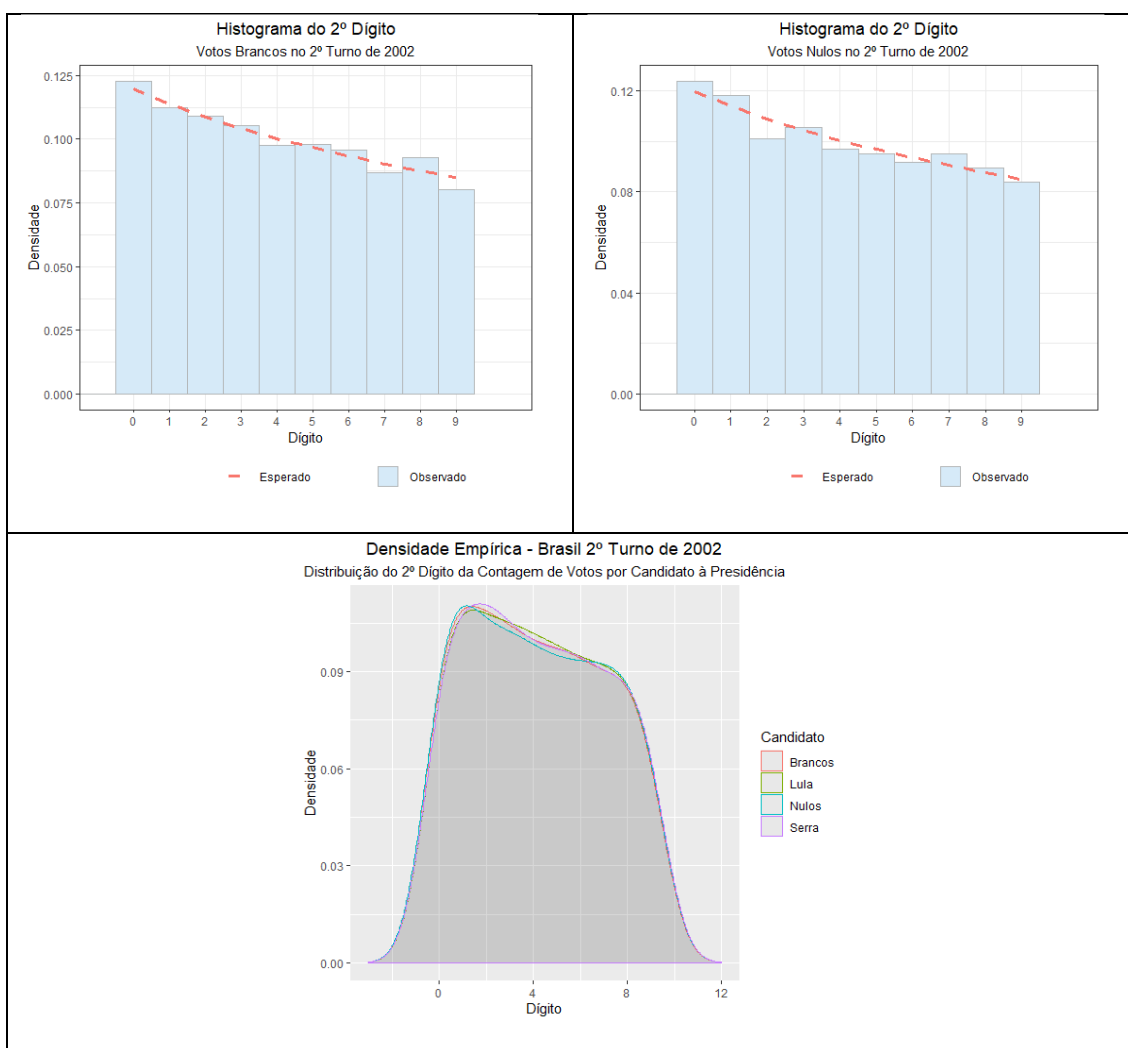


Tabela 20 - Teste Qui-Quadrado - 2002 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística χ^2	gl	p-valor
Lula	4,8278	9	0,8491
Serra	8,9526	9	0,4417
Branco	5,2323	9	0,8136
Nulos	7,4983	9	0,5854

Tabela 21 - Teste G (Teste de Verossimilhança) - 2002 - 2º Turno - Teste 2BL - Dígitos 0 a 9

Candidato	Estatística G	gl	p-valor
Lula	4,8575	9	0,8465
Serra	8,8905	9	0,4474
Branco	5,2396	9	0,8129
Nulos	7,5415	9	0,5809

APÊNDICE 2

```
##### Exemplo de Programa desenvolvido em R para Cálculo
de Aderência à Lei de Benford do 2º dígito para o 1º Turno de 2018
#####

set.seed(25689)

rm(list=ls())

setwd("~/Artigo_eleições")

arquivo <- "votacao_secao_2018_BR.csv"

library(readr)

dados <- read_delim(arquivo,

                    ";", escape_double = FALSE, col_names = TRUE,

                    locale = locale(date_names = "pt", encoding =
"latin1", decimal_mark = ",", grouping_mark = "."), trim_ws = TRUE)

dados <- as.data.frame(dados)

str(as.factor(dados$NR_VOTAVEL))

##### 1º Turno #####
#####

data1t <- dados[dados$NR_TURN0 == 1,]

library(dplyr)

muni <- data1t %>%

  group_by(CD_MUNICIPIO, NR_VOTAVEL) %>%

  summarise(qtde_votos = sum(QT_VOTOS))

muni[,4] <- muni[,3]

names(muni)[4] <- "digito"

temp <- muni$qtde_votos

temp <- substr(temp,2,2)

muni$digito <- temp

muni$digito <- as.numeric(muni$digito)
```



```
Haddad <- muni[muni$NR_VOTAVEL == 13,]
Alckmin <- muni[muni$NR_VOTAVEL == 45,]
Bolsonaro <- muni[muni$NR_VOTAVEL == 17,]
Ciro <- muni[muni$NR_VOTAVEL == 12,]
Amoêdo <- muni[muni$NR_VOTAVEL == 30,]
Branco <- muni[muni$NR_VOTAVEL == 95,]
Nulo <- muni[muni$NR_VOTAVEL == 96,]

digito_haddad <- Haddad$digito[!is.na(Haddad$digito)]
digito_alckmin <- Alckmin$digito[!is.na(Alckmin$digito)]
digito_bolso <- Bolsonaro$digito[!is.na(Bolsonaro$digito)]
digito_ciro <- Ciro$digito[!is.na(Ciro$digito)]
digito_amoedo <- Amoêdo$digito[!is.na(Amoêdo$digito)]
digito_branco <- Branco$digito[!is.na(Branco$digito)]
digito_nulo <- Nulo$digito[!is.na(Nulo$digito)]

digi <- data.frame(c(digito_haddad,
digito_alckmin,digito_bolso,digito_ciro,digito_amoedo,digito_branco,
digito_nulo),

c(rep("Haddad",length(digito_haddad)),rep("Alckmin",length(digito_alck
min)),rep("Bolsonaro",length(digito_bolso)),

rep("Ciro",length(digito_ciro)),rep("Amoêdo",length(digito_amoedo)),re
p("Branco",length(digito_branco)),

rep("Nulos",length(digito_nulo))))

names(digi) <- c("Digito", "Candidato")

#####Vetor com as probabilidades esperadas pela Lei de
Benford de 2º dígito e 1º dígito, respectivamente #####

esperado2 <- c(0.11968, 0.11389, 0.10882, 0.10433, 0.10031, 0.09668,
0.09337, 0.09035, 0.08757, 0.08500)

esperado <- data.frame(A=c(0.11968, 0.11389, 0.10882, 0.10433, 0.10031,
0.09668, 0.09337, 0.09035, 0.08757, 0.08500),

B=seq(-0.5,8.5))
```

```
#####Histogramas e Densidades #####

library(ggplot2)

hist_cand <- function(cand, ano, turno){

  nome <- deparse(substitute(cand))

  f0 <- paste("Votos",nome)

  f1 <- paste(f0, "no")

  f2 <- paste(f1, turno)

  f3 <- paste(f2,"º Turno de", sep="")

  f4 <- paste(f3, ano)

  ggplot(data=cand, aes(cand$digito)) +

    theme_bw() +

    geom_histogram(aes(y=..density.., fill="#D6EAF8"),

                   breaks=seq(-2,9), col="#B3B6B7") +

    labs(title="Histograma do 2º Dígito", subtitle=f4) +

    theme(plot.title = element_text(hjust = 0.5)) +

    theme(plot.subtitle = element_text(hjust = 0.5)) +

    labs(x="Dígito", y="Densidade") +

    geom_line(data=esperado,

              aes(x=esperado$B,y=esperado$A,colour="Esperado"),

              linetype="dashed",size=1.25) +

    scale_fill_manual(values = "#D6EAF8", labels = "Observado") +

    scale_x_discrete(name = "Dígito",

                     limits=seq(-0.5,8.5),

                     labels=c("0","1","2","3","4","5","6","7","8","9")) +

    theme(legend.title = element_text(colour="white", size = 18,
                                       face='bold'),

          legend.key = element_blank(), legend.box = "horizontal",

          legend.position="bottom")

}
```

```

hist_cand(Bolsonaro,2018,1)

hist_cand(Haddad,2018,1)

hist_cand(Ciro,2018,1)

hist_cand(Alckmin,2018,1)

hist_cand(Amoêdo,2018,1)

hist_cand(Branco,2018,1)

hist_cand(Nulo,2018,1)


ggplot(data=digi, aes(Digito)) +

  stat_density(aes(group=Candidato,color = Candidato),bw=0.9115,
alpha=0.05,position="identity")+

  xlim(-3,12) +

  labs(title="Densidade Empírica - Brasil 1º Turno de 2018",
subtitle="Distribuição do 2º Dígito da Contagem de Votos por Candidato
à Presidência") +

  theme(plot.title = element_text(hjust = 0.5)) +

  theme(plot.subtitle = element_text(hjust = 0.5)) +

  labs(x="Dígito", y="Densidade")


freqbolso <- Bolsonaro %>%

  group_by(digito) %>%

  summarise(Freq=n())

freqhad <- Haddad %>%

  group_by(digito) %>%

  summarise(Freq=n())

freqciro <- Ciro %>%

  group_by(digito) %>%

  summarise(Freq=n())

freqalck <- Alckmin %>%

  group_by(digito) %>%

  summarise(Freq=n())

```

```

freqamoedo <- Amoêdo %>%
  group_by(digito) %>%
  summarise(Freq=n())
freqbranco <- Branco %>%
  group_by(digito) %>%
  summarise(Freq=n())
freqnulo <- Nulo %>%
  group_by(digito) %>%
  summarise(Freq=n())

##### Teste Qui-quadrado #####

chisq.test(freqbolso$Freq[1:10], p=esperado2)
chisq.test(freqhad$Freq[1:10], p=esperado2)
chisq.test(freqciro$Freq[1:10], p=esperado2)
chisq.test(freqalck$Freq[1:10], p=esperado2)
chisq.test(freqamoedo$Freq[1:10], p=esperado2)
chisq.test(freqbranco$Freq[1:10], p=esperado2)
chisq.test(freqnulo$Freq[1:10], p=esperado2)

library(DescTools)

GTest(x=freqbolso$Freq[1:10], p=esperado2, correct="none")      #
"none" "williams" "yates"

GTest(x=freqhad$Freq[1:10], p=esperado2, correct="none")
GTest(x=freqciro$Freq[1:10], p=esperado2, correct="none")
GTest(x=freqalck$Freq[1:10], p=esperado2, correct="none")
GTest(x=freqamoedo$Freq[1:10], p=esperado2, correct="none")
GTest(x=freqbranco$Freq[1:10], p=esperado2, correct="none")
GTest(x=freqnulo$Freq[1:10], p=esperado2, correct="none")

```

UM ÍNDICE GLOBAL DE DESEMPENHO DE ATLETAS E EQUIPES DE VOLEIBOL

Mateus Duarte de Menezes Ferreira

estufmg2015@outlook.com

Thiago Rezende dos Santos

thiagords@ufmg.br

Universidade Federal de Minas Gerais

Resumo: Ferramentas estatísticas vêm cada vez mais sendo empregadas para análise de desempenho de equipes e atletas em diferentes esportes. No voleibol, há métricas que mensuram o quão eficazes os jogadores podem ser dentro de quadra nos fundamentos de saque, recepção, levantamento, ataque, bloqueio e defesa exercidos. Um exemplo é a medida de eficiência usada pela Confederação Brasileira de Vôlei, que é a diferença entre os acertos e erros dividido pelo total de ações. Neste trabalho, é introduzido um índice global de eficiência de atletas e equipes baseado na medida de eficiência da CBV em cada um desses fundamentos. Além da estimativa pontual do índice de eficiência, é construído um intervalo de confiança assintótico. Os métodos desenvolvidos são aplicados aos dados da Superliga Brasileira Masculina de voleibol de temporada 2014-2015. Os resultados indicam equilíbrio entre as equipes que disputaram o torneio com um destaque positivo para o Sada Cruzeiro e Sesi e negativo para São Bernardo Vôlei e Voleisul em termos de eficiência. Além disso, destaque para os opositos Wallace e Théó; ponteiros Maurício e Renato; centrais Pedro e Éder; levantadores William e Ricardinho; líberos Serginho e Tiago Brendle, que tiveram as duas maiores eficiências por posição nesta ordem.

Palavras-chave: Eficiência; CBV; Desempenho de atletas e times; Estatística; Proporção; Estimador de máxima verossimilhança; Multinomial

ABSTRACT: Statistical tools are increasingly being employed to analyze the performance of teams and athletes in different sports. In volleyball, some metrics measure how effective players can be on the court in the skills of serve, reception, set, attack, blocking and dig exercised. An example is the efficiency measure used by CBV, which is the difference between successes and errors divided by total actions. In this paper, a global efficiency index of athletes and teams is introduced based on the efficiency measure of the CBV in each of these skills. In addition to the point estimate of the efficiency index, an asymptotic confidence interval is constructed for global efficiency. The methods developed are applied to the data of the Brazilian Superleague Men's volleyball season 2014-2015. The results indicate a balance between the teams that played the tournament with a positive highlight for Sada Cruzeiro and Sesi and negative for São Bernardo Volleyball and Voleisul in terms of efficiency. Besides, stand out for the opposites Wallace and Theo; wing spikers Mauricio and Renato; middle blockers Peter and Éder; setters William and Ricardinho; liberos Serginho and Tiago Brendle, who had the two highest efficiencies per position in this order.

Keywords: Efficiency; CBV; Players and teams performances; Statistics; Proportion; Maximum likelihood estimator; Multinomial

1. INTRODUÇÃO

O aprendizado estatístico dos dados vem sendo cada vez mais explorado em diversas áreas de diferentes esportes, seja desenvolvendo modelos para prever o desempenho de atletas de natação (Hazlewood, 2010), seja no handebol, quando da análise do efeito das posições exercidas pelos jogadores dentro de quadra (Karcher & Buccheit, 2014) ou, ainda, no futebol, através da análise das ações que levam ao sucesso das jogadas (Lago-Ballesteros & Lago-Peñas, 2010) e na realização de simulações para prever os resultados das partidas (Louzada et al., 2015).

No voleibol, por sua vez, alguns estudos já foram publicados visando investigar, explicar, prever e quantificar os fenômenos dessa modalidade, como as variáveis técnicas e táticas que determinam uma vitória ou derrota no set (Rodríguez-Ruiz, 2011), a relação entre a resposta do bloqueador frente ao levantador em bolas de velocidade (Busca & Febrer, 2012) e o efeito de executar um saque ou de recebê-lo de acordo com a configuração de rodízio do time (Laios & Kountouris, 2011).

Um aspecto muito importante a ser destacado é o desempenho que uma equipe ou atleta apresenta dentro de quadra. Por meio desse elemento, pode-se identificar as particularidades de uma equipe ou atleta (Hughes & Franks, 2008) e, assim, otimizar e convergir os treinos para aperfeiçoar as habilidades tanto coletivas como individuais. Nesse sentido, uma das ferramentas que possibilita essa análise cinge-se na utilização de indicadores que mensuram a precisão e a eficiência das habilidades relacionadas a esse esporte (O'Donoghue, 2009). Além disso, a construção de índices de desempenho é de fácil aplicação e já foi utilizada no estudo dos indicadores de rendimento em função do resultado do set (Marcelino et. al, 2010).

A Confederação Brasileira de Voleibol (CBV), órgão responsável pelo voleibol no Brasil, faz uso de uma métrica específica que contabiliza a eficiência dos fundamentos, verificando se a ação executada concluiu-se em sucesso, em continuação ou em erro, calculando-as para times e atletas e ranqueando-os de acordo com seus desempenhos (Marcelino, Mesquita & Afonso, 2008). Em contrapartida, essa classificação utilizada não leva em conta todo o conjunto de ações executadas, sendo apenas por fundamento, inviabilizando-se uma comparação geral entre as equipes onde todas as ações são contabilizadas simultaneamente. Essa limitação também se estende aos atletas e os mesmos não são comparados por posição, apenas por fundamentos isoladamente. É apresentado o melhor atacante, bloqueador e etc, porém não é indicado quem

foi o melhor oposto, ponteiro, líbero e etc em termos de eficiência. Além disso, em geral, só é mostrado um valor de eficiência (uma estimativa) e nenhuma métrica da incerteza dessa estimativa.

O objetivo deste artigo é obter índices/escores globais de eficiência com a medida de eficiência por fundamento adotada pela CBV, considerando todos os fundamentos, simultaneamente, do voleibol para equipes e atletas. Além disso, intervalos de confiança assintóticos para estes índices globais são introduzidos para verificar se há equivalência estatística entre as equipes ou atletas, quando comparados por suas respectivas posições.

Este trabalho está organizado da seguinte forma: na Seção 2, são apresentados os dados e os índices de eficiência da CBV. Na Seção 3, o índice global de eficiência e a construção dos intervalos de confiança assintóticos. Na Seção 4, analisa-se os dados referentes aos times e em sequência os atletas na Seção 5. A conclusão e considerações finais estão na Seção 6.

2. BANCO DE DADOS E A MÉTRICA DA CBV

O banco de dados da Superliga Brasileira Masculina de voleibol foi obtido no site da CBV (superliga.cbv.com.br/14-15/), contendo todas as 22 primeiras rodadas da temporada 2014-2015. Ele é composto pela quantidade de ações classificadas como: sucessos, erros, continuação e pelo total de ações, que é a soma de todas as ações executadas pelos jogadores. Esses dados são disponibilizados por equipes, por atletas e por fundamento, a saber, ataque, bloqueio, defesa, levantamento, recepção e saque.

A CBV possui duas métricas, distinguidas pelo tipo de ação, para definir a eficácia em um determinado fundamento. Para os fundamentos saque e bloqueio, é usada a proporção de acertos. Isto é, a eficiência para o fundamento j da k -ésima equipe/jogador é dada por:

$$e_{jk} = \frac{\# \text{Sucessos}(\text{fundamento}_{jk})}{\# \text{Total de ações}(\text{fundamento}_{jk})}, j = 1(S), 2(B), k = 1, \dots, K;$$

onde K é o número de equipes ou atletas avaliados no campeonato.

Para os fundamentos de recepção, levantamento, ataque e defesa, calcula-se a eficiência através da diferença entre acertos e erros para o fundamento j da k -ésima equipe/jogador:

$$e_{jk} = \frac{\# \text{Sucessos}(\text{fundamento}_{jk}) - \# \text{Erros}(\text{fundamento}_{jk})}{\# \text{Total de ações}(\text{fundamento}_{jk})}, j = 3(D), 4(L), 5(R) \text{ e } 6(A) \text{ e } k = 1, \dots, K.$$

3. MÉTODOS

As comparações serão realizadas da seguinte maneira: primeiro, será computada a eficiência em cada fundamento, para equipes e atletas, de acordo com a métrica adotada pela CBV. Essa métrica de eficiência é adotada neste trabalho porque, além de ser a métrica utilizada pela CBV, ela também é utilizada em outras ligas de voleibol internacionais¹. Em seguida, um índice global, utilizando uma combinação linear dessas eficiências (uma média ponderada), é construído levando em conta os pesos de cada fundamento que são calculados através da proporção de acionamentos frente ao total das ações. Por fim, usando a distribuição amostral da proporção e a distribuição assintótica do estimador de máxima verossimilhança para as proporções do modelo multinomial, constrói-se um intervalo de confiança assintótico para o escore/índice global de eficiência.

3.1 ESCORES PONTUAIS

Inicialmente, para o índice global, é necessário calcular os pesos por fundamento, que é dado pela seguinte fórmula:

$$w_j = \frac{\#Total\ de\ ações\ (fundamento_j)}{\sum_{j=1}^6 \#Total\ de\ ações\ (fundamento_j)}, \text{ para } j = 1, \dots, 6. \quad (1)$$

Para facilitar a notação do índice de eficiência, denomina-se os fundamentos por: 1 – Saque; 2 – Bloqueio; 3 – Defesa; 4 – Levantamento; 5 – Recepção; 6 – Ataque.

Portanto, usando os pesos acima e as eficiências para cada fundamento, o escore global de eficiência S_k da k -ésima equipe ou jogador é dado por:

$$S_k = \sum_{j=1}^6 e_{jk} w_{jk};$$

Onde:

e_{jk} é a eficiência do j -ésimo fundamento para k -ésima equipe ou jogador, com $j = 1, \dots, 6$ e $k = 1, \dots, K$ e K sendo o número de equipes ou atletas avaliados no campeonato.

w_{jk} é o peso do j -ésimo fundamento para k -ésima equipe/jogador, com $j = 1, \dots, 6$ e $k = 1, \dots, K$.

¹ Veja, por exemplo, a liga italiana de voleibol no website
<<http://www.legavolley.it/statistiche/?TipoStat=1.4&lang=en>>.

3.2 ESCORES INTERVALARES

Para a construção dos intervalos de confiança para a eficiência global, é necessário dispor da distribuição amostral dos estimadores da proporção e da eficiência da CBV. Tendo em vista os dados e estatísticas da CBV, podemos usar as medidas de eficiência nos fundamentos.

Para saque e bloqueio, podemos assumir, assintoticamente, sob certas condições, uma distribuição normal, para o estimador da eficiência por se tratar de uma proporção de acertos (Casella & Berger, 2002). Logo, a distribuição da eficiência da k -ésima equipe/jogador para os fundamentos saque e bloqueio ($j = 1, 2$) é dado por:

$$e_{jk} \sim Normal \left(p_{jk}, \frac{p_{jk}(1-p_{jk})}{n_{jk}} \right), \text{ para } j = 1 \text{ e } 2 \text{ e } k = 1, \dots, K; \quad (2)$$

onde n_{jk} é o total de ações da k -ésima equipe ou jogador no j -ésimo fundamento.

Para os fundamentos de ataque, defesa, levantamento e recepção ($j = 3, 4, 5, 6$), considere X_{jk}^A o número de acertos da k -ésima equipe/jogador no fundamento j , X_{jk}^E o número de erros e X_{jk}^C o número de continuações. Reescrevendo a eficiência da CBV, obtém-se:

$$e_{jk} = \frac{\sum_{j=1}^{n_{jk}} X_{jk}^A - \sum_{j=1}^{n_{jk}} X_{jk}^E}{n_{jk}} = \frac{\sum_{j=1}^{n_{jk}} X_{jk}^A - (n_{jk} - \sum_{j=1}^{n_{jk}} X_{jk}^A - \sum_{j=1}^{n_{jk}} X_{jk}^C)}{n_{jk}}, \text{ para } j = 3, 4, 5, 6 \quad (3)$$

e $k = 1, \dots, K$.

Seja $X_{jk} = (X_{jk}^A, X_{jk}^C, X_{jk}^E)$ é um vetor aleatório com três possíveis categorias: acerto, continuação e erro com probabilidades p_{jk}^A, p_{jk}^C , e p_{jk}^E , respectivamente. Para esses fundamentos, pode-se admitir que X_{jk} tem distribuição Multinomial($n_{jk}, p_{jk}^A, p_{jk}^C, p_{jk}^E$) para $j = 3, 4, 5$ e 6 e $k = 1, \dots, K$. Utilizando a distribuição assintótica do estimador de máxima verossimilhança para o vetor de proporções $p = (p_{jk}^A, p_{jk}^C, p_{jk}^E)$ do modelo multinomial e o método delta (Casella & Berger, 2002), obtém-se o seguinte resultado, para o fundamento j da k -ésima equipe/jogador:

$$e_{jk} = \hat{p}_{jk}^A - \hat{p}_{jk}^E \sim Normal \left(p_{jk}^A - p_{jk}^E; \frac{p_{jk}^A(1-p_{jk}^A)}{n_{jk}} + \frac{p_{jk}^E(1-p_{jk}^E)}{n_{jk}} + \frac{2p_{jk}^A p_{jk}^E}{n_{jk}} \right), \quad (4)$$

para $j = 3, 4, 5$ e 6 e $k = 1, \dots, K$.

Para construir a distribuição amostral da eficiência ou escore global levando em conta todos os fundamentos, a hipótese de independência entre os mesmos é assumida e a eficiência global S_k da k -ésima uma equipe ou atleta é dada por:

$$S_k = \sum_{j=1}^6 e_{jk} w_{jk} \sim \text{Normal}(\mu_k, \sigma_k^2), \quad (5)$$

onde:

$$\mu_k = E(S_k) = \sum_{i=1}^6 w_i E(e_{ik});$$

$$\sigma_k^2 = \text{Var}(S_k) = \sum_{i=1}^6 w_i^2 \text{Var}(e_{ik}),$$

sendo e_{ik} calculado de acordo com (2) ou (3), dependendo do fundamento considerado.

Por fim, a distribuição assintótica da eficiência ou escore global em (5) possibilita a construção de um intervalo de confiança assintótico de nível de confiança $1 - \alpha$ para a eficiência global da k -ésima equipe/jogador, cuja forma é dada por:

$$IC_{S_k}^{(1-\alpha)\%} = \left[\hat{\mu}_k \pm |z_{\alpha/2}| \sqrt{\hat{\sigma}_k^2} \right], \text{ para } k = 1, \dots, K, \quad (6)$$

onde $|z_{\alpha/2}| \sqrt{\hat{\sigma}_k^2}$ é chamado de erro padrão da eficiência global e $z_{\alpha/2}$ um quantil de ordem $\alpha/2$ da distribuição normal-padrão.

4. ANÁLISE PARA AS EQUIPES

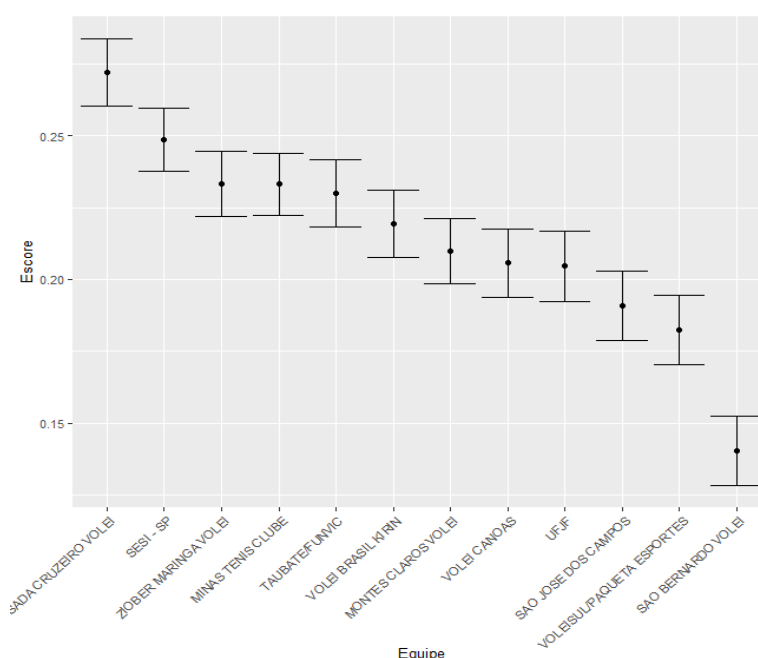
Para os dados das equipes da Superliga masculina brasileira descritos na Seção 2, os pesos obtidos para os fundamentos, de acordo com a Equação (1), foram 0,21 para ataque, 0,10 para bloqueio, 0,14 para defesa, 0,21 para levantamento, 0,15 para recepção, 0,19 para saque. Uma comparação das estimativas dos escores e da classificação das equipes pode ser obtida através da tabela 1 e figura 1. Todas as análises foram conduzidas no software R, versão 3.5.2.

Tabela 1 - Escores globais das equipes e classificação

Equipe	Estimativa Pontual	Estimativa Intervalar
Sada Cruzeiro Vôlei	0,272	[0,260 ; 0,283]
SESI-SP	0,258	[0,237 ; 0,259]
Minas Tênis Clube	0,234	[0,222 ; 0,244]
Ziober Maringá Vôlei	0,233	[0,221 ; 0,243]
Taubaté/Funvic	0,229	[0,218 ; 0,241]
Vôlei Brasil Kirin	0,219	[0,207 ; 0,231]
Montes Claros Vôlei	0,209	[0,198 ; 0,221]
Vôlei Canoas	0,205	[0,193 ; 0,217]
UFJF	0,204	[0,192 ; 0,216]
São José dos Campos	0,190	[0,178 ; 0,202]
Voleisul/Paquetá Esportes	0,182	[0,170 ; 0,194]
São Bernardo Vôlei	0,140	[0,128 ; 0,152]

Fonte: dados da amostra

O Sada Cruzeiro Vôlei é o time que obteve maior escore e o São Bernardo Vôlei ocupa a última posição do ranqueamento. Além disso, nota-se que os escores estão bem próximos, principalmente, dos times que ocupam da terceira a nona posição. Para comparar os clubes, é necessário verificar se os intervalos de confiança estão sobrepostos. Em caso afirmativo, há evidências de que essas equipes têm um desempenho compatível; caso contrário, os desempenhos não são compatíveis.

Figura 1 - Intervalos de confiança de 95% para escores das equipes

Fonte - Autores

A equipe que chega mais próximo do Sada Cruzeiro em termos do escore global de eficiência nos fundamentos é o SESI-SP, porém, com pontuações diferentes, o que é corroborado pelo fato do Sada Cruzeiro ser o campeão deste torneio. Já na parte intermediária, temos sete equipes que são, aproximadamente, equivalentes em termos do escore global, que vão desde o Minas Tênis Clube até a UFJF, indicando um campeonato bastante equilibrado. As duas últimas equipes, Voleisul e São Bernardo foram as que tiveram os menores escores de eficiência global, sendo que esta última tem um índice bem menor das demais equipes e acabou sendo rebaixada para a Superliga B.

5. ANÁLISE PARA OS ATLETAS POR POSIÇÃO

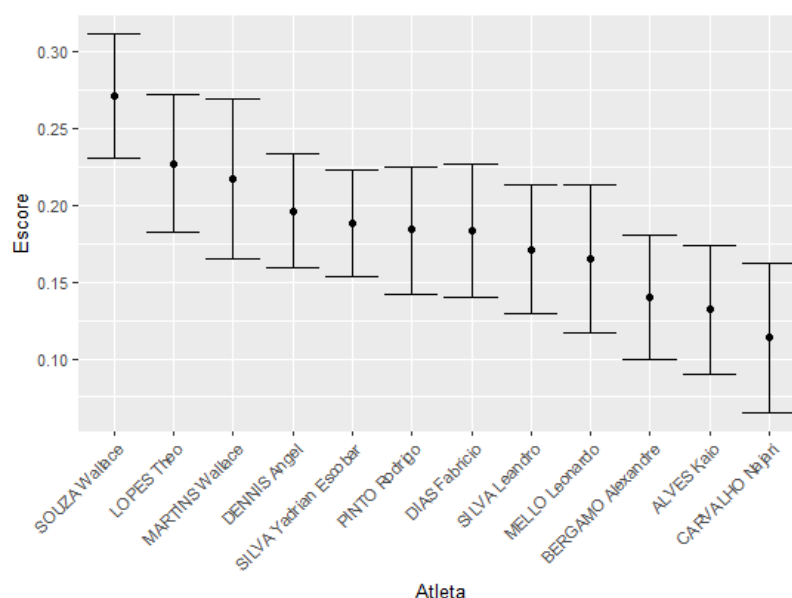
Para os atletas, foi feita a divisão pela posição de cada um, já que alguns jogadores executam mais fundamentos que outros. A titularidade também foi levada em conta, entretanto apenas os resultados de atletas que de fato compõem a equipe principal serão apresentados. Além disso, as composições do escore global de eficiência, em termos dos fundamentos, foram diferentes para determinadas posições. Recepções de opostos e centrais não foram levadas em conta para o cálculo dos pesos. Levantadores não foram pontuados por recepção e ataque e os líberos apenas por defesa, recepção e levantamento. Apenas ponteiros tiveram todas as eficiências dos fundamentos levadas em conta para cálculo do escore global. Como são muitos jogadores por posição,

apresentamos os resultados dos 12 melhores atletas por posição com respeito ao escore global de eficiência, o Top Twelve.

5.1 OPOSTOS

Os pesos dos fundamentos para os opositos foram calculados de acordo com a Equação (1) e são: 0,51 para ataque, 0,11 para bloqueio, 0,13 para defesa, 0,03 para levantamento e 0,22 para saque. O maior peso, para os opositos, é dado para o ataque, visto que é o fundamento mais executado e importante para os mesmos (“virar bola”). As estimativas dos escores globais e seus respectivos intervalos podem ser vistos na figura 2.

Figura 2 - Intervalos de confiança de 95% para escores dos opositos



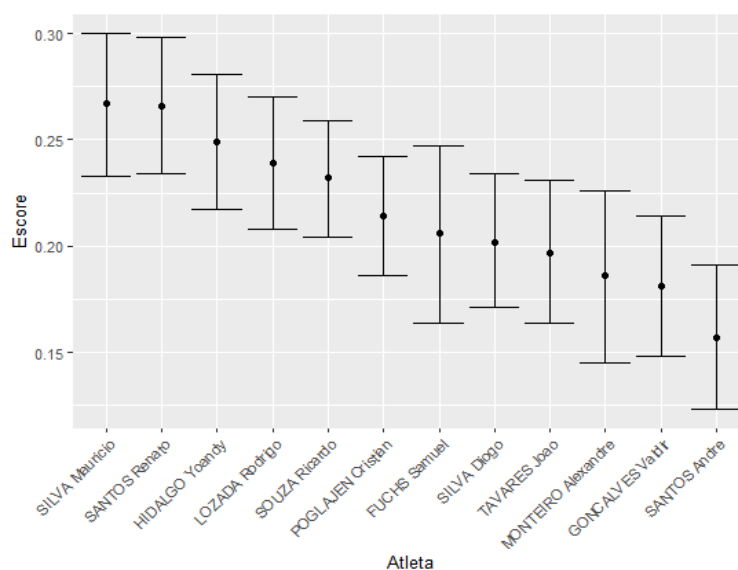
Dentre os opositos titulares, não há muitas disparidades, pois são bastante consistentes no fundamento em que são mais acionados (ataque). Wallace Souza é o atleta com maior escore, porém, apenas outros 3 jogadores podem ser considerados como equivalentes – com eficiências compatíveis. Entretanto, Najari Carvalho, que obteve a menor pontuação, foi compatível com a eficiência de todos os atletas, exceto Wallace (veja, figura 2).

5.2 PONTEIROS

Para os ponteiros, todos os fundamentos foram incluídos para a construção dos escores globais de eficiência. Os ponteiros tiveram pesos dos fundamentos definidos como: 0,27 para ataque, 0,07 para bloqueio, 0,13 para defesa, 0,03 para levantamento, 0,31 para recepção e 0,19 para saque. Os pesos maiores são para os fundamentos de recepção e ataque e é, de fato, o que se espera de um atleta dessa posição. De modo análogo aos opositos,

obtem-se as classificações e intervalos para os ponteiros titulares, selecionando apenas os doze melhores ponteiros em termos do índice global (figura 3).

Figura 3 - Intervalos de confiança de 95% para escores dos ponteiros

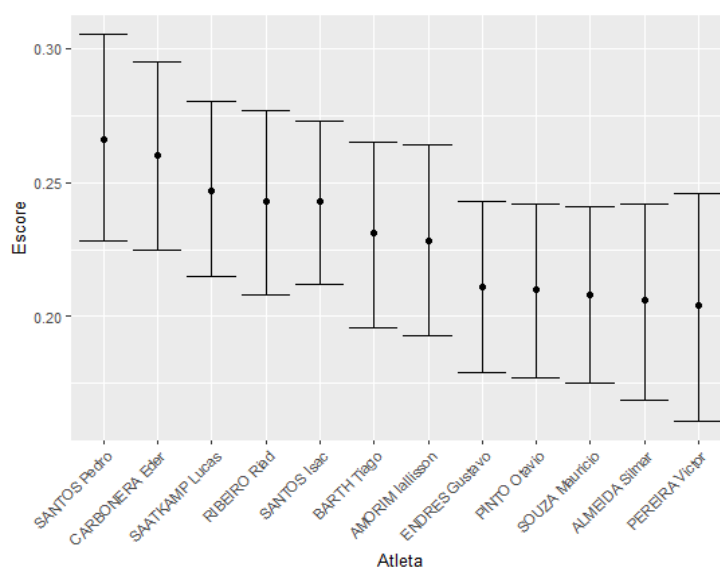


Fonte - Autores

Destacam-se os ponteiros Maurício, Renato e Yoandy Leal que possuem os maiores escores (figura 3). Vale destacar que não fizemos a distinção de ponteiro passador e ponteiro atacante, já que a CBV não considera essa denominação. Os ponteiros são os únicos jogadores que, em teoria, realizam todo tipo de ação dentro de quadra, por isso seus intervalos tendem a ser um pouco mais amplos e os jogadores se equivalem a outros mais facilmente (figura 3). Se um atleta dessa função for muito bom no ataque e ser mediano na recepção provavelmente ele estará no mesmo nível de um outro que seja excepcional na recepção, mas que possui um ataque menos eficiente. Entretanto, para haver consistência de um ponteiro é necessária uma alta eficiência em todos os fundamentos executados, caso contrário os intervalos tendem a ser maiores.

5.3 CENTRAIS

Os pesos dos fundamentos para os centrais foram computados como: 0,27 para ataque, 0,27 para bloqueio, 0,09 para defesa, 0,04 para levantamento e 0,33 para saque. Para os centrais, os maiores pesos são ataque e bloqueio, fundamentos que mais executam. As estimativas dos escores globais e seus respectivos intervalos, para os doze melhores centrais com respeito ao escore global estão na figura 4.

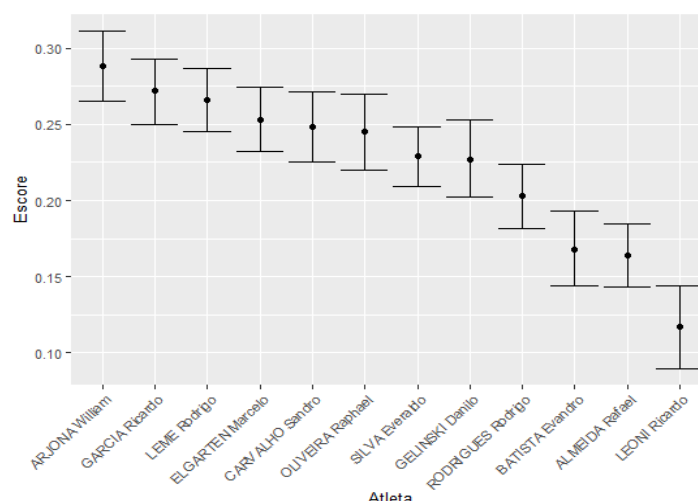
Figura 4 - Intervalos de confiança de 95% para escores dos centrais

Fonte: autores

Sobressaem-se os centrais Pedro e Eder que possuem os maiores escores (figura 4). Os atletas centrais são um tipo diferenciado de jogador, pois realizam jogadas muito rápidas e em formação de bloqueios que, para o desfecho da ação, também dependem de outros fatores, como o atacante adversário ou o bloqueador aliado por exemplo. Muitas vezes, essas ações se resultam em continuações, que não são levadas em conta pela métrica da CBV gerando intervalos maiores; por isso, os intervalos para os atletas centrais são, na maioria das vezes, muito amplos, pois a incerteza sobre as ações que executarão é maior do que se comparado à outras posições. Essa é a única posição em que todos os jogadores podem ser considerados com desempenhos equivalentes.

5.4 LEVANTADORES

Não foram considerados os fundamentos recepção e ataque para os levantadores, já que eles, em geral, raramente exercem essas ações. Para os levantadores, os pesos dos fundamentos foram: 0,72 para levantamento, 0,14 para saque, 0,10 para defesa e 0,04 para bloqueio. Em termos do total de ações executadas por um atleta de uma determinada posição, os pesos refletem nos fundamentos mais executados por esse atleta. Dessa forma, o levantamento tem peso 0,72, que possui o maior peso para jogadores dessa posição. Os escores globais e intervalos de confiança são mostrados na figura 5.

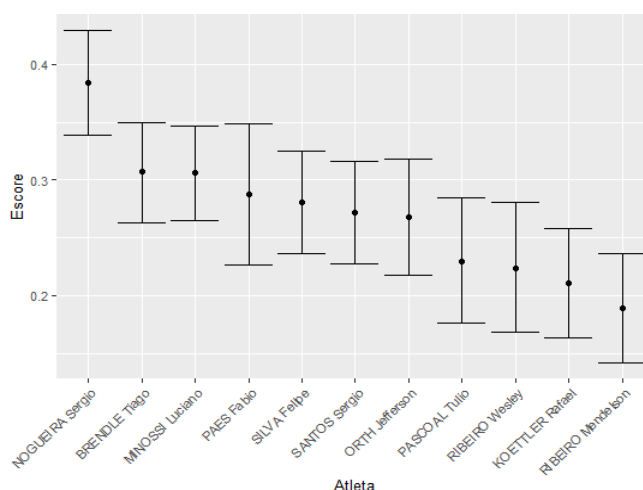
Figura 5 - Intervalos de confiança de 95% para escores dos levantadores

Fonte - Autores

Os levantadores William e Ricardinho tiveram os maiores escores de eficiência global (figura 5). William foi eleito pela CBV o melhor levantador desse campeonato. O levantador realiza a ação intermediária entre passe e ataque sendo responsável por definir qual jogador irá executar esta última. Esse atleta tem um papel importante pois são especializados em executar os levantamentos que irão definir as jogadas ofensivas do time. A diferença dos escores se dá justamente a este fato, que diferencia os jogadores por precisão, consistência e qualidade desse fundamento.

5.5 LÍBEROS

Os líberos executam apenas os fundamentos de recepção, defesa e levantamento, logo, somente, esses foram considerados no escore global de eficiência. Os pesos para os fundamentos foram: 0,48 para recepção, 0,40 para defesa e 0,12 para levantamento, sendo que o primeiro tem o maior peso seguido de perto pela defesa. A estimativa dos escores globais para os líberos e seus respectivos de intervalos de confiança podem ser vistos na figura 6.

Figura 6 - Intervalos de confiança de 95% para escores dos líberos

Fonte: autores

Nota-se que Sergio Nogueira, o primeiro colocado, está a uma distância razoável do segundo colocado, Tiago Brendle (figura 6). O líbero executa sempre apenas dois movimentos fundamentais (toque e manchete) fazendo com que esse tipo de jogador seja extremamente consistente. São atletas que possuem a função mais específica dentro de quadra (receptionar e defender), com isso os intervalos de confiança ficam bem definidos e de tamanhos aproximadamente iguais, diferenciando-os apenas onde está centrado.

6. CONCLUSÃO E CONSIDERAÇÕES FINAIS

Neste trabalho, construímos um índice/escore global de eficiência para atletas e equipes através de uma combinação linear/média ponderada das eficiências nos fundamentos obtidos pela CBV. Os pesos dos fundamentos nessa média foram proporcionais ao total de ações executadas em cada fundamento. Evidentemente, os pesos poderiam ser determinados de outras formas, inclusive atribuindo 1/6 para cada fundamento; porém, optamos por uma forma dos pesos mais intuitiva e de fácil compreensão para não especialistas e profissionais do vôlei. Poderíamos também usar outros métodos para desenvolvimento dos índices de desempenho globais. Um deles, por exemplo, consubstancia-se no uso de Análise de Componentes Principais para a obtenção dos escores pontuais, em conjunto com o método de reamostragem do tipo Bootstrap (Efron e Tibshirani, 1993) para a criação dos intervalos de confiança bootstrap. De fato, chegamos a obter um escore global de eficiência para as equipes, porém os resultados foram muito similares, não havendo diferença no ranqueamento das equipes. Assim, optamos pelo índice/escore global de eficiência baseado em uma média ponderada por ser de fácil interpretação pelos

profissionais do esporte e por sua simplicidade. E, também, temos poucas variáveis/fundamentos, apenas seis.

Diferente da métrica de eficiência da CBV, Marcelino et al. (2010) propõe o uso de uma métrica baseada em pesos para um sucesso, uma continuação ou um erro de uma ação. Nesse método, uma variação dos escores é possível se houver variação dos pesos, que, nesse caso, são diferentes para cada desfecho.

Outro ponto interessante/curioso é a comparação dos escores globais de eficiência com a classificação final do campeonato. De fato, o Sada Cruzeiro, equipe com melhor escore, foi o campeão do torneio e o Sesi-SP, equipe com segundo melhor escore, ficou com o vice. Na parte debaixo da tabela, de acordo com os escores das equipes, o Voleisul teria sido rebaixado junto com o São Bernardo, entretanto, São José dos Campos foi o clube que desceu para a divisão B do torneio, no lugar do Voleisul. Por sua vez, no meio da tabela, houve apenas algumas trocas de, no máximo, uma posição entre os times, se feita a comparação dos índices globais com a classificação final. Nesse sentido, o resultado dos escores das equipes coincidiu, em sua grande maioria, com a classificação final de cada equipe no campeonato.

Em relação aos melhores jogadores, a CBV premia de acordo com melhor ataque, bloqueio, saque, recepção, defesa e melhor levantador, além de premiações específicas como MVP do torneio (*most valuable player*) e “craque da galera”. No campeonato em análise, o prêmio de melhor atacante eleito foi para Wallace Souza (também classificado neste trabalho como melhor oposto), o de melhor bloqueio para Riad Ribeiro (quarto melhor central), o de melhor saque para Luiz, o de melhor recepção para Benedito Canuto (segundo melhor ponteiro classificado), o de melhor defesa para Sergio Nogueira (melhor líbero) e o de melhor levantador para William Arjona (classificado como melhor levantador, de acordo com o índice global de desempenho), desenvolvido neste trabalho.

Quanto ao caso do central Riad, importante destacar que a diferença se deve ao desempenho superior do atleta no fundamento saque; em contrapartida, ele teve aproveitamento inferior aos primeiros colocados, no fundamento ataque, e, por esse motivo, apesar de ser também o melhor bloqueador, não obteve a classificação de melhor central, de acordo com o escore global. Com a construção do índice global de eficiência, podemos, agora, premiar o melhor jogador por posição com um critério de eficácia claro, objetivo e acessível a todos. Isso é bem diferente de premiar atletas só por fundamento sem um critério claro. Registra-se que, em razão da premiação de “craque da galera” ser feita por meio de votação popular pela internet, no presente estudo, não foi feita a

comparação do desempenho do atleta premiado com os índices de desempenho.

É importante lembrar que o MVP do campeonato é o jogador que se destacou em momentos importantes e decisivos para a equipe e, por esse motivo, está sujeito a muita oscilação entre as partidas. Portanto, esse título foi conferido a Yoandry Leal, do Sada Cruzeiro, que foi ranqueado como o terceiro melhor ponteiro e que teve papel decisivo nas fases finais do torneio.

Por fim, um questionamento com respeito ao time que possui melhor desempenho/eficiência pode ser levantado, sugerindo que a equipe de fato vencerá a partida. Entretanto, há outros estudos que levam em conta outros fatores que podem contribuir para a vitória ou derrota de um jogo, como sediar a partida ou jogar como visitante (Alexandros, Panagiotis, & Miltiades, 2012). Neste sentido, o estudo se faz importante para avaliar as necessidades de atletas e equipes em relação aos seus pontos fortes e fracos e, assim, alavancar o desempenho nos fundamentos individuais mais básicos e em questões táticas que envolvem a coletividade.

Trabalhos futuros incluem o aprimoramento dos resultados e técnicas apresentadas, relaxando algumas hipóteses assumidas, bem como a adoção de outras métricas de eficiência que não seja necessariamente a adotada pela CBV e outros pesos para os fundamentos. Seria interessante analisar outras bases de dados de outras ligas de voleibol, incluindo ligas internacionais.

7. REFERÊNCIAS BIBLIOGRÁFICAS

Busca, B., & Febrer, J. (2012). Temporal fight between the middle blocker and the setter in high level volleyball. *International Journal of Medicine and Science of Physical Activity and Sport*, 12(46), 313-327.

Casella, G. & Berger, R. L. (2002). Statistical inference. Pacific Grove, CA: Duxbury.

Heazlewood, T. (2006). Prediction versus reality: The use of mathematical models to predict elite performance in swimming and athletics at the Olympic Games. *Journal of Sports Science and Medicine*, 5(4): 541-547.

Hughes, M., & Franks, I. (2008). *The essentials of performance analysis: an introduction*. London: Routledge.

Karcher, Claude & Buchheit, Martin. (2014). On-Court Demands of Elite Handball, with Special Reference to Playing Positions. *Sports Medicine (Auckland, N.Z.)*. doi: 10.1007/s40279-014-0164-z.

Lago-Ballesteros, J., & Lago-Peñas, C. (2010). Performance in team sports: Identifying the keys to success in soccer. *Journal of Human Kinetics*, 25, 85-91.

Laios, A. & Kountouris, P. (2011). Receiving and serving team efficiency in Volleyball in relation to team rotation. *International Journal of Performance Analysis in Sport*, 11(3), 553-561.

Louzada, F., Suzuki, A. K., Salasar, L. E., Ara, A. & Leite, J. G. (2015). Simulation-based methodology for predicting football match outcomes considering experts' opinions: the 2010 and 2014 football World Cup cases. *Pesquisa Operacional*, 35(3), 577-598.

Marcelino, R., Mesquita, I. & Afonso, J. (2008). The weight of terminal actions in Volleyball. Contributions of the spike, serve and block for the teams' rankings in the World League 2005. *International Journal of Performance Analysis in Sport*, 8(2), 1-7.

Marcelino, R., Mesquita, I., Sampaio, J. & Moraes, J. C. (2010). Study of performance indicators in male volleyball according to the set results. *Brazilian Journal of Physical Education and Sport*, 24(1), 69-78.

Mingoti, S. A. (2005). Análise de dados através de métodos de estatística multivariada: uma abordagem aplicada. Belo Horizonte: Editora UFMG. 1ª edição.

O'Donoghue, P. (2009). Interacting performances theory. *International Journal of Performance Analysis in Sport*, 9(1), 26-46.

R Development Core Team. (2005). *R: A language and environment for statistical computing*, R Foundation for Statistical Computing, Vienna, Austria.

Rodriguez-Ruiz, D., Quiroga, M. E., Miralles, J. A., Sarmiento, S., de Saá, Y. & García-Manso, J. M. (2011). Study of the technical and tactical variables determining set win or loss in top-level European men's volleyball. *Journal of Quantitative Analysis in Sports*, 7(1), 1-15.

OS EFEITOS DA CRISE FINANCEIRA AMERICANA DE 2008 NA PRODUÇÃO FÍSICA INDUSTRIAL BRASILEIRA: UMA ANÁLISE DE INTERVENÇÃO

Koki Fernando Oikawa

kfoikawa@gmail.com

UniCapital

Resumo: O objetivo desse artigo consiste em avaliar se a crise financeira americana de 2008 exerceu impacto real sobre volume de produção brasileira. Em geral, os artigos com essa temática ou são fundamentalmente descritivos ou utilizam modelos VAR. Assim, será utilizada uma abordagem que não é tão frequentemente utilizada, a saber, a Análise de Intervenção. Para tal, será necessário estimar um modelo SARIMA para, em seguida, realizar uma análise de intervenção para a série temporal de produção física industrial (PFI), do IBGE, de modo a testar se houve, de fato, impacto real na série temporal da PFI. Os resultados sugerem que houve impacto real e negativo e que a recuperação da atividade econômica ocorreu de forma gradual, porém, proeminente.

Palavras-chave: Séries Temporais; ARIMA; SARIMA; Análise de Intervenção; Crise Bancária

Abstract: This article aimed to assess whether 2008 US financial crisis had a real impact on brazilian production volume index. In general, articles with this theme are either fundamentally descriptive or use VAR models. Thus, an approach not so often used, namely Intervention Analysis, will be applied. For this, it will be necessary to estimate a SARIMA model to perform an Intervention Analysis for the IBGE industrial physical production (IPP) time series in order to test whether there was, in fact, real impact on the IPP time series. The results suggest that there was a real and negative impact and the economic activity recovery occurred gradually, but prominently.

Keywords: Time series; ARIMA; SARIMA; Intervention Analysis; Banking crisis

1. INTRODUÇÃO

A crise financeira americana de 2008 reverberou nas economias internacionais. Do ponto de vista doméstico, o PIB americano (dados do FMI/IFS disponíveis no IPEA-Data) evoluiu, em 2007, em 4,87%. No ano da crise, chegou, ainda, a crescer, mas, em patamares de 1,87%. Em 2009 houve retração (-2,22%) e apenas em 2010 voltou a crescer 3,76%.

Na Europa, o PIB da Zona do Euro (dados do FMI/IFS disponíveis no IPEA-Data, em US\$) crescia a uma taxa de 5,70% em 2007 e, no ano da crise, apresentou desempenho de 2,39%. Em 2009 houve queda de 2,71%, voltando a crescer somente em 2010, a uma taxa de 2,77%. Terazi e Senel (2011) fez um estudo detalhado sobre os efeitos da crise bancária americana sobre a economia dos países da Europa Central e do Leste Europeu (República Tcheca, Estônia, Letônia, Lituânia, Hungria, Polônia, Eslovênia, Eslováquia, Romênia e Bulgária). De acordo com os autores, os impactos nas economias desses países foram diferentes, por razões, circunstâncias distintas. Apenas a título de ilustração, o impacto na Hungria foi muito forte, também, em função de políticas errôneas do seu próprio passado. Eslovênia, Eslováquia e República Tcheca mostraram-se em uma posição, talvez, um pouco melhor e a Bulgária teria sido o único país em uma situação relativamente favorável. As análises dos impactos em cada um dos países levaram em consideração 5 indicadores macroeconômicos: i) Crescimento real do PIB, ii) taxa de desemprego, iii) Arrecadação fiscal dos governos (em % do PIB), iv) Gastos públicos (em % do PIB) e iv) Empréstimos consolidados aos governos (% do PIB).

Com relação aos impactos na economia mundial, a variação real anual do PIB mundial (dados do FMI/IFS disponíveis no IPEA-Data), em 2007 e 2008 foram respectivamente de 5,4% e 2,76%, refletindo uma queda na taxa de crescimento de praticamente 50%. Em 2009, houve uma queda real (crescimento negativo) de 0,61%, voltando a crescer 5,27% em 2010.

Vários analistas passaram a ressaltar características específicas da crise financeira americana. Bresser-Pereira (2009) destacou o fato de ser uma crise bancária e não de balanço de pagamentos, bastante comum em países em desenvolvimento. Krugman (2009) analisa em detalhes a dinâmica de formação dessa crise como uma combinação de certos fatores, dentre eles, a formação de uma bolha no mercado de ações. Junior e Filho (2008) chamam atenção para a formação de uma ampla liquidez nos EUA decorrente da condução da própria política monetária adotada pelo Banco Central americano desde 2001 a qual produziu um forte movimento de valorização dos ativos imobiliários com indícios de bolha especulativa.

Em razão de certas peculiaridades que são características dessa crise surge, então, o problema de avaliar se ela produziu (ou não) impactos reais na economia brasileira e, em caso positivo, se esses impactos foram temporários ou permanentes.

A análise de intervenção surgiu em um artigo seminal de Box e Tiao (1975). Nesse artigo, os autores fizeram uso da recém-desenvolvida metodologia Box-Jenkins de séries temporais que permitia fazer inferência a partir de séries univariadas. Por meio da estimação de um modelo ARIMA Sazonal (SARIMA) combinado com certas funções de transferência, seria possível fazer esse tipo de análise, isto é, avaliar se determinado evento em algum(ns) ponto(s) do tempo teria(m) produzido efeitos sobre a trajetória de uma série e se esses efeitos seriam temporários ou permanentes.

Nesse artigo original, Box e Tiao fizeram duas aplicações em problemas econômicos: a) Dados de ozônio em Los Angeles. Nesse caso o objetivo consistia em avaliar o impacto de duas intervenções agindo no sentido de redução da camada de ozônio. Em 1960 foi inaugurada a auto-estrada Golde State (intervenção 1) e em 1966 foram criadas novas regulações para mudança no funcionamento dos motores dos carros (intervenção 2). Os resultados para a intervenção 1 foram evidentes enquanto que para a intervenção 2, a redução no ozônio ocorreu de forma foi gradual, sendo muito pequena nos períodos seguintes. b) Taxa de mudança no índice de preços ao consumidor americano. Em agosto de 1971 foram instituídos alguns controles, de forma que o objetivo consistia em avaliar a influência daqueles nos índices de preços. Também foram consideradas duas intervenções: a fase I é a fase de implementação dos controles enquanto que a fase II seriam os efeitos percebidos nos índices, nos períodos subsequentes. Os resultados mostraram que os controles produziram efeitos reais na queda dos índices para a fase I enquanto que, na fase II, os efeitos mostraram-se incertos.

No Brasil, a técnica de análise de intervenção tem sido utilizada, em Economia, principalmente em temas que procuram avaliar o impacto de planos econômicos em determinadas séries. Fava, Vera. L. e Denisard C. O. Alves (1997) em “Indicador de Movimentação Econômica, Plano Real e Análise de Intervenção” procuraram avaliar o impacto do Plano real no indicador Imec (média ponderada de dez variáveis que estivessem relacionadas com a movimentação econômica. Dentre elas, estavam: número de passageiros em ônibus urbanos, passageiros em vôos nacionais, consultas ao SPC etc). Os autores concluíram que o Plano Real produziu impactos reais e de longo prazo sobre o indicador Imec.

Na área agrícola, Maura, M. D. Santiago, Maria de Lourdes B. Camargo e Mario Antonio Margarido (1996) em “Detecção e Análise de Outliers em Séries Temporais de Índices de Preços Agrícolas no Estado de São Paulo”, utilizaram, também, uma técnica de análise de intervenção baseada na técnica de detecção de outliers para investigar a influência do processo inflacionário no período 1966-94 sobre os preços agrícolas. Essa técnica é utilizada geralmente quando não é possível identificar precisamente o(s) período(s) da(s) intervenção(ões).

Tupy (2015) chama a atenção para o fato de que na literatura econômica, a maioria dos trabalhos que visam avaliar o impacto de choques monetários sobre variáveis econômicas têm utilizado a metodologia VAR. Dentre eles, Tupy (2015) destaca Carlino e DeFina (1996), Rodriguez-Fuentez e Padrón-Marrero (2008), Ciccarelli et al. (2013), Fraser et al. (2012) e Silva (2012). Tupy (2015) cita, ainda, Bueno (2008) para o qual a aplicação da técnica VAR visa a obtenção de respostas para a trajetória de uma série a partir da ocorrência de um choque estrutural de uma outra série.

Esse artigo tem por finalidade a análise do impacto da crise bancária americana sobre a economia brasileira com um enfoque diferente, a saber, a partir da aplicação da técnica de Análise de Intervenção sobre a série temporal de Produção Física Industrial – PFI – do IBGE. Assim, a resposta da produção física industrial não será o foco da análise, mas, outro sim, se houve impacto real sobre a economia brasileira e, em caso positivo, obter uma medida da velocidade de recuperação. Para a realização da análise de intervenção, será necessário a estimação de um modelo SARIMA (ARIMA sazonal). Em todas as análises que se seguiram, contou-se com o apoio do software estatístico R.

Esse artigo encontra-se organizado da seguinte maneira: na seção 2 serão feitas algumas considerações sobre a crise financeira americana, bem como a determinação do período-chave para a data da intervenção para a devida análise. Na seção 3, serão apresentados os fundamentos teóricos de forma a estimar a avaliar a adequação de um modelo SARIMA para uma série temporal. Na seção 4, as principais funções de intervenção e de transferência serão apresentadas, de modo a fundamentar a análise de intervenção. Os resultados serão apresentados na seção 5 e as conclusões, na seção 6.

2. A CRISE FINANCEIRA AMERICANA

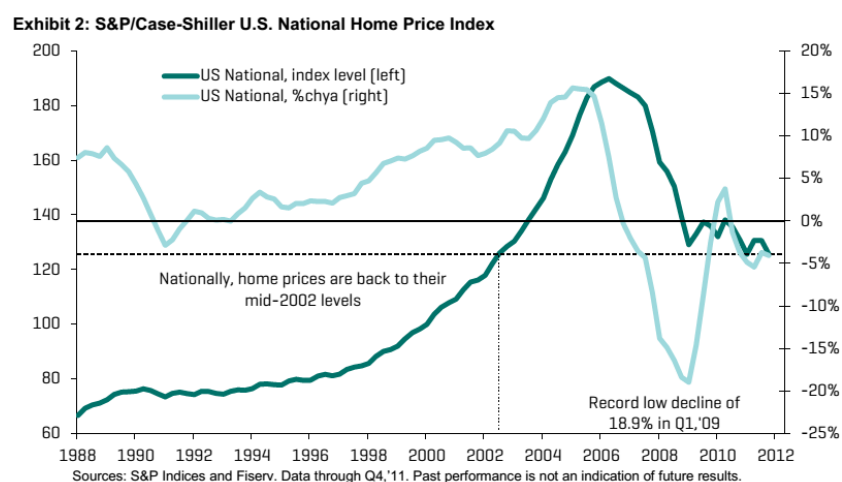
A crise financeira americana originou-se de uma crise em seu mercado imobiliário local, precisamente no mercado de hipotecas. Júnior e Filho (2008) destacam alguns fatores. No contexto do mercado imobiliário, taxas relativamente baixas de juros vinham sendo praticadas pelo Federal Reserve desde o início da década passada, logo após os atentados terroristas (11 de setembro de 2001). Isso acabou por favorecer a expansão do crédito, em geral, e da modalidade de crédito imobiliário, em particular. Assim, surgiram estímulos para que as famílias americanas se endividassem mais, contratando novas linhas de crédito e ou fazendo novas hipotecas. Esses estímulos também se estenderam para famílias e agentes ou com históricos ruins de crédito (baixa renda, sem garantias, dívidas anteriores não saldadas etc), que formariam o volume de crédito denominado *subprime* que seriam, em outras palavras, os empréstimos com grandes chances de inadimplência.

A combinação desses fatores mencionados (baixas taxas de juros, expansão do crédito imobiliário induzido, em larga medida, por subprimers), promoveu, de 1997 a 2006, uma forte valorização dos preços dos imóveis no mercado americano chegando, praticamente, a triplicar o valor. Esse forte crescimento pode ser observado no gráfico da Figura 1, elaborado pela S&P/Case-Shiller, em seu relatório “Home Price Indices: 2011 Year in Review”, disponível em seu site.

O gráfico da Figura 1 chama-nos, ainda, a atenção para o comportamento simultâneo de duas curvas:

- a) a do índice de preços dos imóveis nos EUA, e
- b) o retorno anual do próprio índice, isto é, a taxa de variação percentual do índice dentro de um período de 12 meses.

De acordo com esse gráfico da Figura 1, é possível notar que no instante em que as duas curvas se cruzam pela primeira vez, precisamente em 2006, a curva dos preços dos imóveis até continuam crescendo, mas, já em um ritmo bem mais lento e já próximo do ponto de máximo.

Figura 1 – Índice de retorno e de preços dos imóveis dos EUA de 1988 a 2012

É especificamente nesse instante que o retorno, isto é, a taxa de variação percentual do índice (dentro de um período de 12 meses) entra praticamente em queda livre até o início de 2009 (seu ponto mais crítico). Mas, em meados de 2008, ou seja, um pouco antes desse ponto ser alcançado, a crise do setor é, então, deflagrada, transmitindo enorme instabilidade nos mercados internacionais.

A amplificação dessa crise imobiliária no mercado americano acabou produzindo impactos nos mercados globais, em larga medida, pela forma como foram arquitetadas as operações financeiras que deram suporte ao crescimento do próprio mercado imobiliário: as hipotecas subprime foram transformadas em títulos negociáveis nos mercados financeiros internacionais, ou seja, em títulos securitizados.

Por outro lado, é necessário, ainda, destacar que, a partir de 2004, o Federal Reserve passou a imprimir uma política monetária mais restritiva que se estendeu até 2006, elevando, segundo Lemos e Bittencourt (2007), a taxa de juros de 1% ao ano para 5,25% ao ano. Assim, refinar, ou até mesmo arcar com as prestações de um financiamento imobiliário ficava mais difícil para os agentes subprimers, ocasionando inadimplência em massa, uma vez que as taxas de juros nesses tipos de contratos eram pós-fixadas.

Mas, enquanto o Federal Reserve adotava uma política monetária mais restritiva, Junior e Filho (2008) mostram, em números, que de 2002 a 2006, as hipotecas subprime securitizadas cresceram de US\$ 121 bilhões para US\$ 483 bilhões respectivamente e, no mesmo período, a proporção dessas hipotecas securitizadas, que estava em torno de 52,7%, subiu para 80,5%. Isso significa que, além de o volume de empréstimo ter crescido enormemente (para quase meio trilhão de dólares) e de ter sido direcionado a um público com históricos

duvidosos, a grande maioria seria transformada em títulos securitizados para serem negociados nos mercados financeiros internacionais. Além dos fatores da origem da crise já destacados anteriormente por Junior e Filho (2008), Cechin e Montoya (2017) ressaltam que a concorrência entre os agentes financeiros atuantes no mercado de hipotecas também contribuiu para a propagação de diversos tipos de contratos, uma vez que agiam no intuito de atrair tomadores de maior risco. Daí, o efeito devastador que a inadimplência poderia exercer. Cechin e Montoya (2017) relatam ainda que essas carteiras de crédito foram securitizadas e transformadas em ativos financeiros (CDO) para serem revendidos.

Já em março de 2008, a falência do 5º maior banco de investimentos dos EUA (Bear Stearns) fez com que governo americano concedesse uma linha de crédito em torno de US\$ 30 bilhões ao JP Morgan para que esse pudesse fazer a aquisição.

Em julho, as empresas de crédito imobiliário Fannie Mae e Freddie Mac (ambas detêm quase a metade do mercado de hipotecas nos EUA) apresentaram sérios problemas de liquidez.

Em setembro, para não fecharem suas portas, o governo americano optou por estatizar essas duas empresas, adquirindo 80% das ações preferenciais de ambas as empresas.

Também em setembro, o Lehman Brothers acabou entrando com pedido de concordata na Corte de Falências de Nova York. A decisão do governo americano de não oferecer apoio financeiro a esse banco aprofundou a crise, deixando os mercados financeiros internacionais em pânico.

Logo em seguida (ainda em setembro), a maior companhia de seguros dos EUA, a AIG, anuncia problemas de liquidez sendo que, desta vez, o governo americano optou por conceder empréstimos de US\$ 85 bilhões. Além disso, em 15 de setembro ocorreu a maior falência da história americana (Lehman Brothers). Em função desse retrospecto, e para os propósitos desse trabalho, o mês de setembro será considerado como a data de intervenção.

Não podemos deixar de destacar o fato de que vários analistas chamaram a atenção para a (provável) existência de bolhas especulativas. Krugman (2009) chama a atenção para a bolha especulativa no mercado de ações que teria se formado na década de 90, antecedendo a bolha do mercado imobiliário que estaria por vir. Essas duas bolhas, segundo o autor, poderiam ser percebidas no comportamento de duas curvas: i) o índice de preços dos imóveis dos EUA (também baseado naquele elaborado pela S&P/Case-Shiller)

e ii) a relação preço-lucro (P/L), construído por Robert Shiller da Universidade de Yale.

Krugman (2009) observou (apesar de apontar para a necessidade de certos cuidados com a análise baseada nesses dois índices) que enquanto o índice de preços dos imóveis manteve-se relativamente estável na década de 90 (com indícios de crescimento somente no final da década), a relação preço-lucro, no mesmo período, cresceu de forma extremamente acelerada. Essa bolha no mercado de ações teria ocorrido em razão de dois motivos que favoreceram as condições de sua formação: i) extremo otimismo acerca de lucros potenciais teria recebido muita atenção e ii) a sensação crescente de uma economia segura ou, em outras palavras, a crença de que grandes recessões seriam coisas do passado. De acordo, ainda, com Krugman (2009), esses dois motivos teriam contribuído para pressionar os preços das ações para cima, isto é, favoreceram a formação de uma bolha no mercado de ações.

Especificamente no Brasil, Costanzi (2009), citada em Cechin e Montoya (2017), descreve que os efeitos da crise bancária americana ocorreram, basicamente, por meio de três canais: i) Forte contração do crédito como resultado do aumento da aversão ao risco, ii) Deterioração das expectativas dos empresários resultando na queda dos investimentos das indústrias e iii) Recessão econômica mundial, com consequências para as exportações brasileiras (queda) e para os preços dos bens exportáveis (queda). Outros estudos também traziam análises com resultados semelhantes. Freitas (2009) destaca o conservadorismo dos bancos brasileiros (aversão ao risco) como responsável pela contração do crédito na economia e, conseqüentemente, redução na atividade econômica.

3. O MODELO SARIMA

Antes de abordarmos a modelagem SARIMA, devemos introduzir inicialmente a modelagem ARIMA, uma vez que a primeira também constitui um modelo ARIMA, mas com uma característica específica, a saber, a de ser sazonal.

Os procedimentos de identificação, estimação e diagnóstico de um modelo ARIMA de série temporal, também conhecidos como “Metodologia de Box-Jenkins”, serão descritos nessa seção. Em seguida, a estratégia de ajuste de um modelo SARIMA será apresentada.

3.1 MODELOS ARIMA (p,d,q)

Quando nos referimos a uma família de modelos ARIMA(p,d,q) em geral, significa dizer que qualquer modelo ARIMA pode ser decomposto em três

partes: AR (auto-regressivo), I (integrado) e MA (médias móveis). Uma série pode ser descrita como um processo auto-regressivo de ordem “p”, integrado de ordem “d” e ou de médias móveis de ordem “q”. Iniciaremos pelo processo AR(p).

3.1.1 PROCESSO AR(p)

Dizemos que $X_t, t \in \mathbb{Z}$ é um processo auto-regressivo de ordem p, $X_t \sim \text{AR}(p)$ se o processo puder ser descrito por

$$X_t = \phi_0 + \phi_1 X_{t-1} + \dots + \phi_p X_{t-p} + \varepsilon_t, \quad (3.1)$$

onde $\phi_0, \phi_1, \dots, \phi_p$ são parâmetros reais, $\varepsilon_t, t \in \mathbb{Z}$, é ruído branco com $\mathbb{E}(\varepsilon_t) = 0$ e $\text{Var}(\varepsilon_t) = \sigma_\varepsilon^2$. Portanto, em um processo $\text{AR}(p)$, X_t depende dos últimos p valores do processo.

Considerando o operador defasagem B, é comum também escrevermos as observações passadas, por exemplo

$$B^0 X_t = X_t, \quad B^k X_t = X_{t-k}, \quad k \neq 0.$$

Desse modo, um processo AR(p) da forma (3.1) também pode ser escrito como

$$\phi(B)X_t = \phi_0 + \varepsilon_t,$$

$$\text{onde } \phi(B) = 1 - \phi_1 B - \phi_1 B^2 - \dots - \phi_p B^p.$$

Sua esperança é dada por

$$\begin{aligned} \mathbb{E}(X_t) &= \mathbb{E}(\phi_0 + \phi_1 X_{t-1} + \dots + \phi_p X_{t-p} + \varepsilon_t) \\ &= \mathbb{E}(\phi_0) + \mathbb{E}(\phi_1 X_{t-1}) + \dots + \mathbb{E}(\phi_p X_{t-p}) + \mathbb{E}(\varepsilon_t) \\ &= \phi_0 + \phi_1 \mathbb{E}(X_{t-1}) + \dots + \phi_p \mathbb{E}(X_{t-p}) \\ &= \frac{\phi_0}{1 - \phi_1 - \dots - \phi_p}, \end{aligned}$$

sendo que a expressão de $\mathbb{E}(X_t)$ só será válida se X_t for estacionário. Em termos mais técnicos, X_t será estacionário se, e somente se, $\phi(B) = 1 - \phi_1 B - \phi_1 B^2 - \dots - \phi_p B^p$ não possuir raízes no círculo unitário, ou seja, se $\phi(B) \neq 0$, para todo $B \in \mathbb{C}$, tal que $|B| = 1$.

3.1.2 PROCESSO MA (q)

Dizemos que $X_t, t \in \mathbb{Z}$ segue um processo de médias móveis de ordem q, se for possível escrever

$$X_t = \theta_0 + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_p \varepsilon_{t-p}, \quad (3.2)$$

onde $\theta_0, \theta_1, \dots, \theta_p$ são parâmetros reais e ε_t é ruído branco com $\mathbb{E}(\varepsilon_t) = 0$ e $\text{Var}(\varepsilon_t) = \sigma_\varepsilon^2$, com a constante θ_0 eventualmente igual a zero.

Diferentemente de um processo auto-regressivo, um processo de médias móveis caracteriza-se por constituir uma combinação de erros presente e passado.

De modo similar, a equação (3.2) também pode ser reescrita na forma $X_t = \theta(B)\varepsilon_t$, sendo que $\theta(B) = 1 - \theta_1 B - \dots - \theta_q B^q$.

A esperança de um processo $MA(q)$ é dado por

$$\begin{aligned}\mathbb{E}(X_t) &= \mathbb{E}(\theta_0 + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q}) \\ &= \mathbb{E}(\theta_0) + \mathbb{E}(\varepsilon_t) - \mathbb{E}(\theta_1 \varepsilon_{t-1}) - \dots - \mathbb{E}(\theta_q \varepsilon_{t-q}) \\ &= \theta_0 + \mathbb{E}(\varepsilon_t) - \theta_1 \mathbb{E}(\varepsilon_{t-1}) - \dots - \theta_q \mathbb{E}(\varepsilon_{t-q}) \\ &= \theta_0.\end{aligned}$$

3.1.3 PROCESSO ARMA (p,q)

Para fins de simplificação, consideraremos ϕ_0 e θ_0 iguais a zero. Mas é perfeitamente possível trabalhar sem essa simplificação.

Um processo $X_t, t \in Z$ é descrito por um processo $ARMA(p, q)$ se puder ser escrito na forma

$$X_t = \phi_0 + \phi_1 X_{t-1} + \dots + \phi_p X_{t-p} + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q}, \quad (3.3)$$

na qual $\phi_0, \phi_1, \dots, \phi_p, \theta_1, \dots, \theta_p$ são parâmetros reais, ε_t é ruído branco com $\mathbb{E}(\varepsilon_t) = 0$ e $Var(\varepsilon_t) = \sigma_\varepsilon^2$.

Nesse caso, a equação (3.3) também pode ser reescrita na forma $\phi(B)X_t = \phi_0 + \theta(B)\varepsilon_t$.

A esperança de um processo $ARMA(p, q)$, supondo estacionariedade para o processo X_t , é dado por

$$\begin{aligned}\mathbb{E}(X_t) &= \mathbb{E}(\phi_0 + \phi_1 X_{t-1} + \dots + \phi_p X_{t-p} + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q}) \\ &= \phi_0 + \mathbb{E}(\phi_1 X_{t-1}) + \dots + \mathbb{E}(\phi_p X_{t-p}) - \mathbb{E}(\theta_1 \varepsilon_{t-1}) - \dots - \mathbb{E}(\theta_q \varepsilon_{t-q}) \\ &= \phi_0 + \phi_1 \mathbb{E}(X_{t-1}) + \dots + \phi_p \mathbb{E}(X_{t-p}) - \theta_1 \mathbb{E}(\varepsilon_{t-1}) - \dots - \theta_q \mathbb{E}(\varepsilon_{t-q}) \\ &= \phi_0 + \mathbb{E}(X_t) (\phi_1 + \dots + \phi_p) \\ \mathbb{E}(X_t) - \mathbb{E}(X_t) (\phi_1 + \dots + \phi_p) &= \phi_0 \\ \mathbb{E}(X_t) (1 - \phi_1 - \dots - \phi_p) &= \phi_0 \\ \mathbb{E}(X_t) &= \phi_0.\end{aligned}$$

3.1.4 PROCESSO ARIMA (p,d,q)

Box, Jenkins e Reinsel (1994) chamam atenção para o fato de que séries econômicas frequentemente apresentam um comportamento não-estacionário, mas ainda capaz de ser escrito por um modelo ARMA (p,q). Neusser (2016) aponta que transformações nesse tipo de séries são necessárias para torná-las estacionárias e que nem sempre somente a primeira diferença é suficiente. Nesse caso, se, em uma série não-estacionária gerada por um processo ARMA (p,q), foi necessário aplicar um certo número d de diferenças no intuito de torná-la estacionária, dizemos, então, que essa série é integrada de ordem d, ou seja, $X_t \sim I(d)$, cujo processo é descrito por um ARIMA(p,d,q) com expressão

$$\phi(B)(1 - B)^d X_t = \theta(B)\varepsilon_t,$$

se $\phi_0 = 0$.

Em relação ao valor “d” referente à integração “I” de um modelo ARIMA(p,d,q), o mesmo pode ser zero.

3.1.5 PROCESSO SARIMA (p,d,q)

Muitas séries temporais econômicas podem eventualmente exibir um certo padrão sazonal. Box, Jenkins e Reinsel (1994) especificam que uma série temporal apresenta um comportamento com certa periodicidade s quando algumas similaridades tendem a ocorrer após s intervalos de tempo (trimestral, semestral, anual). A classe de modelos SARIMA(p,d,q)x(P,D,Q)s constitui uma extensão da classe de modelos ARIMA(p,d,q) com expressão dada por

$$\phi(B)\Phi(B^s)(1 - B)^d(1 - B^s)^D X_t = \theta(B)\Theta(B^s)\varepsilon_t,$$

com a introdução de um operador de diferença sazonal, $(1 - B^s)^D$, e mais dois polinômios em B^s , $\Phi(B^s)$ e $\Theta(B^s)$, com graus P (ordem autoregressiva sazonal) e Q (ordem de médias móveis sazonal), respectivamente. No caso, os operadores $\Phi(B^s)$ e $\Theta(B^s)$ são dados por

$$\Phi(B^s) = 1 - \Phi_1 B^s - \Phi_2 B^{2s} - \dots - \Phi_P B^{Ps}$$

e

$$\Theta(B^s) = 1 - \theta_1 B^s - \theta_2 B^{2s} - \dots - \theta_Q B^{Qs}.$$

A seguir, será apresentado um teste estatístico, de modo a verificar a estacionariedade de uma série.

3.2 IDENTIFICAÇÃO

Nessa seção, introduziremos dois métodos que auxiliam na identificação dos valores p e q de um modelo $ARMA(p, q)$: o primeiro consiste nas funções (e seus gráficos) de autocorrelação (fac) e de autocorrelação parcial (facp). O segundo método trata do Critério de Informação de Akaike (AIC).

3.2.1 FUNÇÕES DE AUTOCORRELAÇÃO (FAC) E DE AUTOCORRELAÇÃO PARCIAL (FACP)

A identificação dos valores de (p, d, q) passam por uma análise gráfica das funções estimadas de autocorrelação (fac) ρ_j e de autocorrelação parcial (facp) ρ_{jj} .

O estimador da função de autocorrelação (fac), em um contexto de um processo estacionário, é dado por

$$r_j = \frac{c_j}{c_0}, \quad j = 0, 1, \dots, T-1,$$

onde c_j é um estimador para a função de auto-covariância (facv) γ_j e

$$c_j = \frac{1}{T} \sum_{t=1}^{T-j} [(X_t - \bar{X})(X_{t+j} - \bar{X})], \quad j = 0, 1, \dots, T-1$$

e $\bar{X}$ a média amostral.

Morettin (2006) ressalta que a fac é mais apropriada para identificar modelos de médias móveis MA. Para o caso de modelos AR, a facp tem se mostrado mais útil. Um método capaz de encontrar a facp para um processo estacionário com fac ρ_k baseia-se nas equações de Yule-Walker. Para isso temos,

$$\rho_k = \phi_{k1}\rho_{j-1} + \phi_{k2}\rho_{j-2} + \dots + \phi_{kk}\rho_{j-k}, \quad j = 1, \dots, k.$$

A partir da qual obtemos as equações

$$\begin{bmatrix} 1 & \rho_1 & \rho_2 & \dots & \rho_{k-1} \\ \rho_1 & 1 & \rho_2 & \dots & \rho_{k-2} \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ \rho_{k-1} & \rho_{k-2} & \rho_{k-3} & \dots & 1 \end{bmatrix} \begin{bmatrix} \phi_{k1} \\ \phi_{k2} \\ \vdots \\ \phi_{kk} \end{bmatrix} = \begin{bmatrix} \rho_1 \\ \rho_2 \\ \vdots \\ \rho_k \end{bmatrix}.$$

Seja ϕ_{kj} o j -ésimo coeficiente de um modelo auto-regressivo de ordem k , ou seja de um modelo $AR(k)$, de tal forma que ϕ_{kk} seja seu último coeficiente. Solucionando sucessivamente para $k = 1, 2, 3, \dots$, obteremos ϕ_{kk} como segue:

$$\phi_{11} = \rho_1$$

$$\phi_{22} = \frac{\begin{vmatrix} 1 & \rho_1 \\ \rho_1 & \rho_2 \end{vmatrix}}{\begin{vmatrix} 1 & \rho_1 \\ \rho_1 & 1 \end{vmatrix}} \rightarrow \phi_{22} = \frac{\rho_2 - \rho_1^2}{1 - \rho_1^2}$$

$$\phi_{33} = \frac{\begin{vmatrix} 1 & \rho_1 & \rho_2 \\ \rho_1 & 1 & \rho_2 \\ \rho_2 & \rho_1 & \rho_3 \end{vmatrix}}{\begin{vmatrix} 1 & \rho_1 & \rho_2 \\ \rho_1 & 1 & \rho_1 \\ \rho_2 & \rho_1 & 1 \end{vmatrix}} \rightarrow \phi_{33} = \frac{\rho_3 + \rho_2^2 \rho_1 + \rho_1^3 - 2\rho_1 \rho_2 - \rho_1^2 \rho_3}{1 - 2\rho_1^2 - \rho_2^2}.$$

Em geral, teremos

$$\phi_{kk} = \frac{|P_k^*|}{|P_k|},$$

onde P_k é a matriz de autocorrelação e P_k^* é a matriz P_k com a última coluna substituída pelo vetor de autocorrelação. Os valores de ϕ_{kk} encontrados em função de k são conhecidas como função de autocorrelação parcial (facp).

Por meio da obtenção dos valores de fac e facp, Morettin (2006) apresenta resumidamente o seguinte procedimento para identificar modelos ARIMA:

Um processo AR(p) tem facp $\phi_{kk} \neq 0$, para $k \leq p$ e $\phi_{kk} = 0$, para $k > p$;

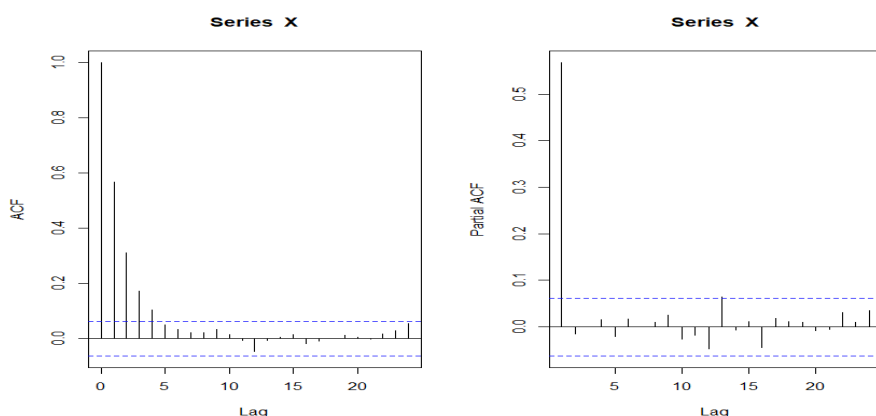
Um processo MA(q) tem facp que se comporta de maneira similar à fac de um processo AR(p): é denominada por exponenciais e/ou senóides amortecidas;

Um processo ARMA(p, q) tem facp que se comporta como a facp de um processo MA puro.

De modo a ilustrar o procedimento de identificação acima foram feitas três simulações no software R para as três possibilidades: AR(1), MA(1) e ARMA(1).

Conforme o procedimento, um processo AR(1) é identificado por apresentar fac dominada por exponenciais (decrece exponencialmente) enquanto que na facp ocorre um truncamento no lag p , conforme Figura 2. Os parâmetros da simulação foram: i) $\alpha = 0,6$, ii) $n=1.000$ e iii) semente (123).

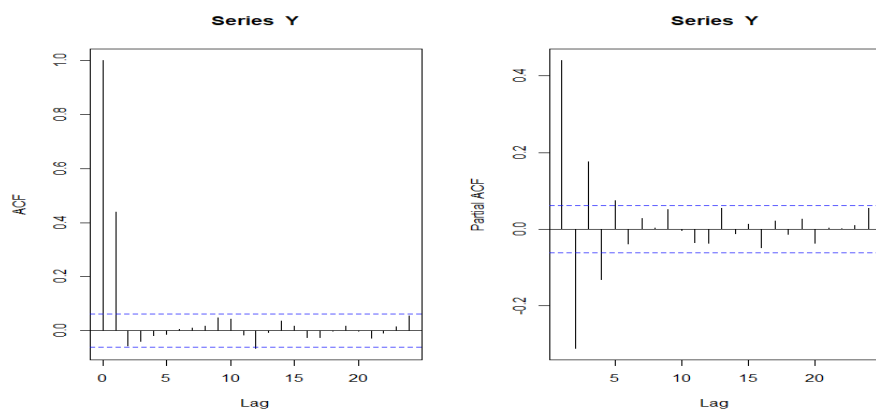
Figura 2 – PROCESSO AR(1)



Fonte - O autor

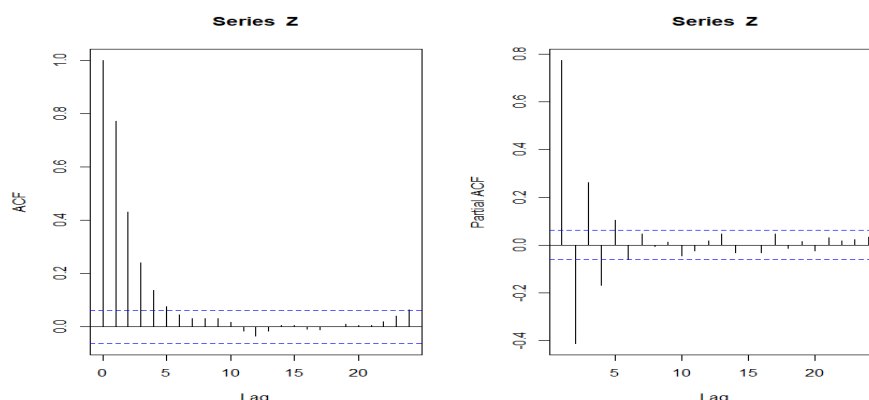
No caso de um processo MA, o método de identificação é exatamente o contrário: A fac apresentará um truncamento no lag q , enquanto que a facp exibirá um decrescimento exponencial, conforme Figura 3 a seguir. Uma observação importante: o gráfico da fac inicia sempre pelo lag zero, que possui valor igual a 1 (lembre que no caso em que $j = 0$, $r_0 = \frac{c_0}{c_0} = 1$) e deve, portanto, ser sempre desconsiderado. Os parâmetros da simulação foram: i) $ma = 0,7$, ii) $n=1.000$ e iii) semente (123).

Figura 3 – PROCESSO MA(1)



Fonte - O autor

Já um processo ARMA, sua facp é similar à facp de um processo MA puro, isto é, deve exibir um comportamento de decrescimento exponencial, como na Figura 4.

Figura 4 – PROCESSO ARMA(1,1)

Fonte - O autor

É importante notar que para um processo ARMA, o procedimento nada diz sobre sua fac, o que acaba dificultando um pouco mais sua identificação. Os parâmetros da simulação foram: i) $ar = 0,6$ e $ma = 0,7$, ii) $n=1.000$ e iii) semente (123).

3.2.2 CRITÉRIO DA INFORMAÇÃO DE AKAIKE - AIC

Uma forma alternativa de identificar a(s) ordem(ens) de modelos $ARMA(p, q)$ consiste em analisar o Critério AIC. Esse critério permite identificar um modelo cujas ordens k e l minimizam o valor seu critério:

$$AIC(k, l) = \ln \hat{\sigma}_{k,l}^2 + \frac{2(k + l)}{T},$$

onde $\hat{\sigma}_{k,l}^2$ é o estimador de máxima verossimilhança de σ^2 para um modelo $ARMA(k, l)$.

Assim, melhor será o modelo que apresentar o menor valor do critério AIC.

Antes de avançarmos, é necessário destacarmos informações importantes sobre o comportamento dos gráficos da fac, quando os dados referem-se a uma série com sazonalidade.

Quando a série é sazonal, essa constatação poderá ser observada nos próprios gráficos de fac e facp, a partir dos quais será possível perceber um certo padrão. Esse padrão manifesta-se na forma de picos altos que se repetem de forma sazonal, ou seja, com uma certa periodicidade (trimestral, semestral, anual etc). Tal padrão é comum de ser encontrado em dados de vendas, arrecadação de impostos, produção, taxas de desemprego entre outras. Maiores detalhes sobre os métodos de ajuste de um modelo SARIMA estão na seção 3.5.

3.4 ESTIMAÇÃO

Com relação ao processo de estimação, tanto para um modelo AR(p), quanto para os modelos MA(q) e ou ARMA(p,q), os fundamentos são basicamente os mesmos. Introduziremos, a seguir, apenas a título de ilustração, os procedimentos de um modelo ARMA(p,q), já supondo estacionariedade por parte dos dados.

Considere um modelo $ARMA(p, q)$ com $p + q + 1$ parâmetros pertencentes ao vetor paramétrico $\xi = (\phi, \theta, \sigma_\varepsilon^2)$, no qual $\phi = (\phi_1, \dots, \phi_p)$, $\theta = (\theta_1, \dots, \theta_q)$ e σ_ε^2 , a variância do erro. Dessa forma,

$$X_t = \phi_1 X_{t-1} + \dots + \phi_p X_{t-p} + \varepsilon_t - \theta_1 \varepsilon_{t-1} - \dots - \theta_q \varepsilon_{t-q}.$$

Podemos reescrever esse modelo isolando ε_t , ou seja,

$$\varepsilon_t = X_t - \phi_1 X_{t-1} - \dots - \phi_p X_{t-p} + \theta_1 \varepsilon_{t-1} + \dots + \theta_q \varepsilon_{t-q}.$$

Supondo normalidade em ε_t , e que sejam dados p valores para X_t e q valores para ε_t , denotados por $\mathbf{X}_t^*$ e $\boldsymbol{\varepsilon}_t^*$ respectivamente, então a função de verossimilhança será dada por

$$L(\xi | \mathbf{X}, \mathbf{X}_t^*, \boldsymbol{\varepsilon}_t^*) = (2\pi)^{-n/2} (\sigma_\varepsilon)^{-n} \exp \left\{ \frac{1}{2\sigma_\varepsilon^2} \sum_{t=1}^n (X_t - \phi_1 X_{t-1} - \dots - \phi_p X_{t-p} + \theta_1 \varepsilon_{t-1} + \dots + \theta_q \varepsilon_{t-q})^2 \right\}.$$

Tomando o logaritmo de L e considerando o vetor $\boldsymbol{\eta} = (\phi, \theta)$, teremos

$$l(\xi | \mathbf{X}, \mathbf{X}_t^*, \boldsymbol{\varepsilon}_t^*) \propto -n \log \sigma_\varepsilon - \frac{S(\boldsymbol{\eta} | \mathbf{X}, \mathbf{X}_t^*, \boldsymbol{\varepsilon}_t^*)}{2\sigma_\varepsilon^2},$$

onde

$$S(\boldsymbol{\eta} | \mathbf{X}, \mathbf{X}_t^*, \boldsymbol{\varepsilon}_t^*) = \sum_{t=1}^n \varepsilon_t^2(\boldsymbol{\eta}, \mathbf{X}, \mathbf{X}_t^*, \boldsymbol{\varepsilon}_t^*),$$

que é chamada soma de quadrados (SQ) condicional.

Desse modo, maximizar $l(\xi | \mathbf{X}, \mathbf{X}_t^*, \boldsymbol{\varepsilon}_t^*)$ equivale a minimizar $S(\boldsymbol{\eta} | \mathbf{X}, \mathbf{X}_t^*, \boldsymbol{\varepsilon}_t^*)$ e os estimadores de MV serão estimadores de MQ.

3.5 DIAGNÓSTICO

Uma vez determinado o modelo ARIMA(p, d, q) que melhor se ajusta ao conjunto de dados, dá-se início a etapa de análise de diagnóstico do modelo em questão. Em outras palavras, a análise de diagnóstico implica em um teste para verificar a existência (ou não) de autocorrelação nos resíduos.

Ljung-Box (1978) propuseram uma versão melhorada de um teste anterior, de forma que as hipóteses envolvidas são dadas por:

$$\begin{cases} H_0: \text{os resíduos são não autocorrelacionados} \\ H_1: \text{os resíduos são autocorrelacionados} \end{cases}.$$

Com relação à estatística do teste de autocorrelação residual, temos que

$$Q(K) = \frac{1}{n(n+2)} \sum_{k=1}^h \frac{\hat{\rho}_k^2}{n-k},$$

na qual n é o tamanho da amostra, $\hat{\rho}_k^2$ é a estimativa da autocorrelação no lag k , e a estatística $Q(K)$ possuindo distribuição χ^2 com $K = h - p - q$ graus de liberdade. A hipótese nula, H_0 , é rejeitada quando $Q(K) > \chi_{1-\alpha, K}^2$ e o objetivo geral desse teste consiste em avaliar a existência de autocorrelação residual.

3.6 MODELO SARIMA

A ideia central da modelagem SARIMA consiste, basicamente, em uma vez detectada sazonalidade, aplicar mais uma vez a metodologia ARIMA na base de dados. Desse modo, é comum representar um modelo SARIMA na forma

$$ARIMA(p, d, q) \times (P, D, Q)_s.$$

Nesse caso, P e Q referem-se, respectivamente, às ordens auto-regressiva e de médias móveis da sazonalidade, enquanto que D refere-se ao número de diferenças efetuadas na série no intuito de se remover o efeito da sazonalidade e s , a magnitude da sazonalidade (trimestral, semestral, anual etc).

Dessa forma, ao detectar sazonalidade, é recomendável tomar uma diferença na série como um todo, de ordem correspondente à sazonalidade percebida. Assim, se, por exemplo, for detectada sazonalidade anual, deve-se tomar uma diferença de ordem 12 ($D = 1, s = 12$), em toda a série, de forma a remover o efeito dessa sazonalidade.

A partir dessa nova base de dados é necessário obter os novos gráficos de *fac* e *facp*.

Shumway and Stoffer (2011), apresentaram uma tabela (Tabela 3.3 da página 155), a qual auxilia na identificação de um processo *ARIMA* sazonal, ou *SARIMA*. O método de identificação proposto pelos autores para um processo *SARIMA* pode ser resumido da seguinte forma:

- Se a fac é decrescente nos lags ks , $k = 1, 2, \dots$ e a facp “trunca” após o lag P_s , então o processo seria $AR(P)_s$, ou seja, auto-regressivo de ordem P e de sazonalidade s ;
- Se a fac “trunca” após o lag Qs e a facp é decrescente nos lags ks , $k = 1, 2, \dots$, então o processo seria $MA(Q)_s$, ou seja, de médias móveis de ordem Q e de sazonalidade s ;
- Se a fac é decrescente nos lags ks , $k = 1, 2, \dots$ e a facp é decrescente nos lags ks , $k = 1, 2, \dots$, então o processo seria $ARMA(P, Q)_s$, ou seja, auto-regressivo de ordem P , também de médias móveis de ordem Q e de sazonalidade s .

É importante frisar, também, que os primeiros lags (1,2,3,...) tanto da fac, quanto da facp são importantes para determinar a ordem dos componentes *não sazonais* p e q . Nesse caso, Shumway e Stoffer (2011) consideram duas possibilidades:

- Se, nas funções fac e facp, os primeiros lags forem estatisticamente significantes e decrescentes, então, $p > 0$ e $q > 0$, de forma que a sugestão é $p = 1$ e $q = 1$;
- Se for considerado, por outro lado, que a facp “trunca” no lag p , então, essa observação deverá ser levada em conta e, simultaneamente, $q = 0$.

Após a estimação de um modelo SARIMA, assim como em um modelo ARIMA, os resíduos (do modelo) também deverão ser analisados (teste diagnóstico) para que seja possível avançar para uma análise de intervenção.

4. ANÁLISE DE INTERVENÇÃO

A análise de intervenção constitui uma técnica de séries temporais que tem por finalidade avaliar o impacto que uma determinada intervenção (mudança na legislação, inovação tecnológica, modificação estrutural, etc) pode exercer na trajetória de uma série temporal. Dessa forma, por intervenção, podemos considerar o acontecimento de algum evento, relevante para o problema e num dado instante do tempo T , capaz de influenciar, de alguma forma, a evolução da série em questão. As intervenções, em geral, podem ser descritas por meio de certas funções específicas.

4.1 FUNÇÕES DE INTERVENÇÃO

No tocante a sua manifestação, uma intervenção pode se originar de forma abrupta ou gradual. Com relação a sua duração, ela pode produzir efeitos temporários ou permanentes. Existe, também, a possibilidade de alguns tipos de combinações entre suas formas e duração. Para maiores detalhes, Box e Tiao (1975) apresenta um quadro contendo esses arranjos.

Fundamentalmente, as intervenções são descritas por dois tipos de funções indicadoras (binárias):

a) Função impulso (*pulse function*)

$$X_{j,t} = I_t^T = \begin{cases} 0, & t \neq T, \\ 1, & t = T; \end{cases} \quad (4.1)$$

b) Função degrau (*step function*)

$$X_{j,t} = S_t^T = \begin{cases} 0, & t < T, \\ 1, & t \geq T. \end{cases} \quad (4.2)$$

A *função impulso* é utilizada quando o efeito da intervenção é *temporário*. Já a *função degrau* retrata uma intervenção que produz efeitos *permanentes*.

Para cada tipo de intervenção, é possível designar uma função de transferência Z_t apropriada. Se estivermos interessados em uma classe mais geral dessas funções de transferências, considerando, inclusive, a possibilidade de múltiplas intervenções ao longo do tempo, essa classe mais geral de funções de funções de transferências é comumente escrita na forma

$$Z_t = \sum_{j=1}^k v_j(B) X_{j,t} + N_t, \quad (4.3)$$

na qual:

- $v_j(B), j = 1, 2, \dots, k$, expressam funções racionais na forma $\frac{\omega_j(B)B^{b_j}}{\delta_j(B)}$, de forma que $\omega_j(B) = \omega_{j,0} - \omega_{j,1}B - \omega_{j,2}B^2 - \dots - \omega_{j,s}B^s$ e $\delta_j(B) = 1 - \delta_{j,1}B - \dots - \delta_{j,r}B^r$ são polinômios em B e b_j é o tempo de defasagem para o início do efeito da *j-ésima* intervenção;
- $X_{j,t}, j = 1, 2, \dots, k$ representam as variáveis indicadoras de intervenção do tipo impulso ou degrau como descritos em (4.1) e (4.2);

- N_t é a série temporal isenta do efeito das intervenções e denomina-se série residual.

A expressão (4.3) expressa o fato de que *antes* da intervenção, $Z_t = N_t$, uma vez que a função de intervenção $X_{j,t}$, seja ela qual for, é igual a zero. É importante ressaltar que deve ser ajustado um modelo SARIMA para a série residual N_t .

4.2 FUNÇÕES DE TRANSFERÊNCIA

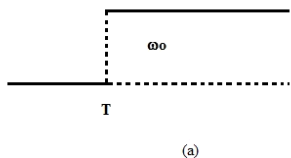
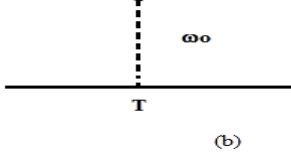
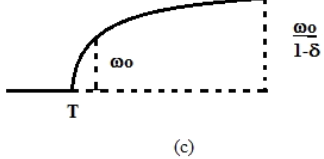
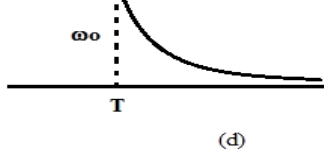
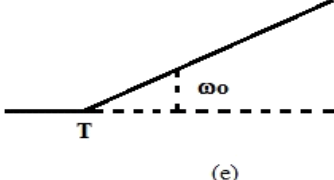
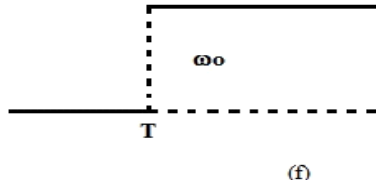
Com relação à função de transferência citada em (4.3), para cada $X_{j,t}$ é possível atribuir uma função $v_j(B)$ distinta, de modo a caracterizar uma função de transferência Z_t específica.

A composição de uma função $v_j(B)$ com uma função de intervenção $X_{j,t}$, recebe o nome de estrutura da função de transferência. A Tabela 1 reúne algumas das funções de transferências mais comumente utilizadas e encontram-se em Morettin e Toloi (2006).

De acordo com a Tabela 4.1 acima, temos as seguintes situações:

- a) Se a função de transferência $v(B)X_t = \begin{cases} 0, & t < T \\ \omega_0, & t \geq T \end{cases}$, a intervenção será abrupta com efeitos permanentes;
- b) Se $v(B)X_t = \begin{cases} 0, & t \neq T \\ \omega_0, & t = T \end{cases}$, então, a intervenção também será abrupta mas com efeitos apenas temporários, isto é, apenas em $T = t$;
- c) Se $v(B)X_t = \begin{cases} 0, & t < T \\ \frac{\omega_0}{1-\delta B}, & t \geq T \end{cases}$, a intervenção ocorrerá gradualmente, com início em $T = t$, com efeitos permanentes;
- d) Por outro lado, se $v(B)X_t = \begin{cases} 0, & t \neq T \\ \frac{\omega_0}{1-\delta B}, & t = T \end{cases}$, a intervenção será abrupta em $T = t$, porém, com efeitos temporários;
- e) Se $v(B)X_t = \begin{cases} 0, & t < T \\ \frac{\omega_0}{1-B}, & t \geq T \end{cases}$, a intervenção em $T = t$ será gradual com efeitos crescentes e permanentes na série;
- f) Se $v(B)X_t = \begin{cases} 0, & t \neq T \\ \frac{\omega_0}{1-B}, & t = T \end{cases}$, ocorrerá a mesma situação descrita no item (a), isto é, intervenção abrupta com efeitos permanentes;

Tabela 1 – Estrutura da Função de Transferência

$v(B)$	$X_t = \begin{cases} 0, & \text{se } t < T \\ 1, & \text{se } t \geq T \end{cases}$	$X_t = \begin{cases} 0, & \text{se } t \neq T \\ 1, & \text{se } t = T \end{cases}$
ω_0	 (a)	 (b)
$\frac{\omega_0}{1 - \delta B}$ $ \delta < 1$	 (c)	 (d)
$\frac{\omega_0}{1 - B}$	 (e)	 (f)

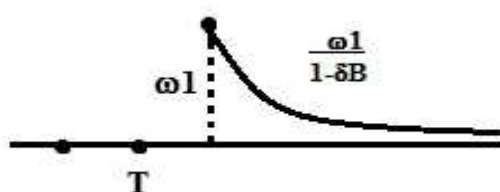
Fonte - Morettin, P. A. e Toloí, C. M. C. (2006),

Existem algumas variantes das funções de transferência. Uma variante da função descrita no item d da Tabela 1 é extremamente importante para os objetivos desse trabalho.

Box e Tiao (1975) apresentaram uma função de transferência, a qual, não incorpora os efeitos da intervenção imediatamente. Essa função de transferência seria apropriada para captar os efeitos da intervenção apenas no período seguinte, ou seja, apenas em $T = t + 1$. No artigo dos autores supracitados, a Figura B(d), da página 72 expressa, especificamente, essa característica. Essa função de transferência é dada por

$$v(B)X_t = \begin{cases} 0, & t \neq T \\ \frac{\omega_1 B}{1 - \delta B}, & t = T \end{cases},$$

a qual, graficamente, encontra-se disposta na Figura 5.

Figura 5 – Função de Transferência com um Período de Defasagem

Fonte - Box and Tiao (1975)

Devemos notar que nesse tipo de função de transferência, os efeitos da intervenção não seriam sentidos no momento da intervenção, ou seja, em $T = t$. Temos, então, que ω_0 não aparecerá na função de transferência, uma vez que a suposição é que não há efeitos a serem medidos no instante da intervenção ou, em outras palavras, que $\omega_0 = 0$.

Nessa mesma linha de raciocínio, vamos propor, pelos motivos que serão elucidados na próxima seção, uma função de transferência com características parecidas, isto é, com intervenção abrupta e efeitos temporários, mas, com defasagem de ordem 2. Essa função de transferência será dada pela expressão (4.4) abaixo:

$$v(B)X_t = \begin{cases} 0, & t \neq T \\ \frac{\omega_2 B}{1 - \delta B}, & t = T \end{cases} \quad (4.4)$$

É importante ressaltar que, no caso da função de transferência dada por (4.4), apesar de a intervenção ter ocorrido em $T = t$, os efeitos dessa intervenção só seriam sentidos em $T = t + 2$, ou seja, ω_0 e ω_1 não seriam, nesse caso, estatisticamente significantes.

5. RESULTADOS

Com o objetivo de avaliar os efeitos da crise financeira americana na produção brasileira, optou-se pelo índice de Produção Física Industrial – Brasil (PFI) do IBGE, construído a partir dos dados primários da Pesquisa Industrial Mensal de Produção Física (PIM-PF).

Sua amostragem é intencional, abrange 944 produtos e 6.200 empresas e representa 62% do valor da produção do Censo Industrial de 1985. Os dados foram obtidos do Banco de Dados Agregados – SIDRA – do IBGE, diretamente do seu site.

Os dados estão compreendidos entre janeiro de 2008 e dezembro de 2011. Conforme mencionado no final da seção 2, o mês de setembro de 2008 é

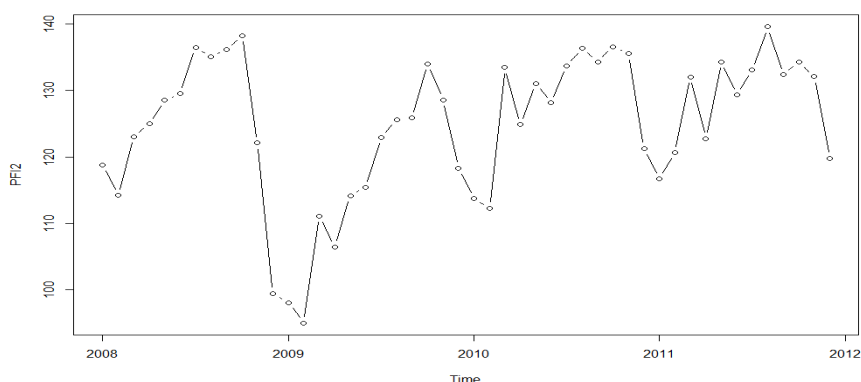
o período adotado para o mês da deflagração da crise financeira americana ou, em outras palavras, o momento da intervenção na série de PFI.

No gráfico da Figura 6 é possível perceber que, mesmo que a crise tenha sido deflagrada em setembro de 2008 (9ª obs.), o índice PFI continuou a subir por mais um mês para, somente em novembro de 2008, dois meses depois, começar a sofrer os efeitos da crise. Houve, portanto, uma defasagem de ordem 2, constituindo o motivo principal da escolha da função de transferência descrita em (4.4).

Na base de dados considerados, o período da intervenção (setembro de 2008) é a 9ª observação que pode ser devidamente identificada no gráfico da Figura 6.

Ao realizarmos um teste de raiz unitária (Dickey-Fuller Aumentado) para a série, não se rejeitou a hipótese nula, ao nível de significância de 0,05 (p-valor de 0,1298). Logo não se rejeitou a hipótese de não estacionariedade da série, ao nível de significância indicado.

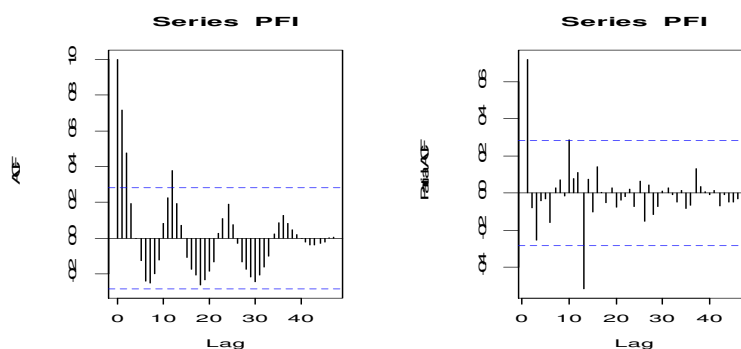
Figura 6 – Gráfico em linha da série PFI-Brasil



Fonte - O autor.

Uma vez que se trata de um índice de produção, os gráficos de fac e facp também são bastante elucidativos para averiguar a possibilidade de sazonalidade na série. Nesse intuito, seguem ambos os gráficos da série PFI na Figura 7.

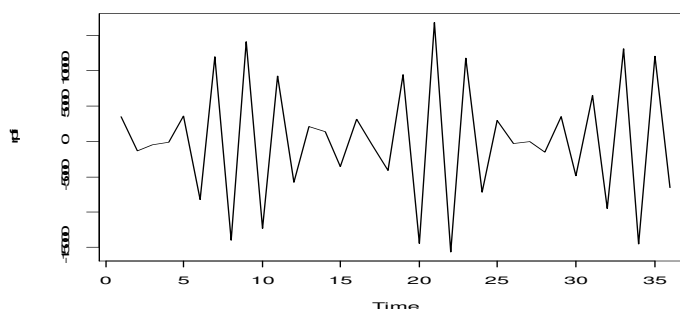
Figura 7 – FAC E FACP DA SÉRIE PFI - BRASIL



Fonte - O autor.

Com base na Figura 7, é possível perceber um certo padrão de sazonalidade na série que, aparentemente, parece se repetir de doze em doze meses. Logo, é necessário tomar uma diferença na série de ordem 12, de modo a remover o efeito da sazonalidade. A Figura 8 apresenta o gráfico dessa nova série, agora aparentemente estacionária.

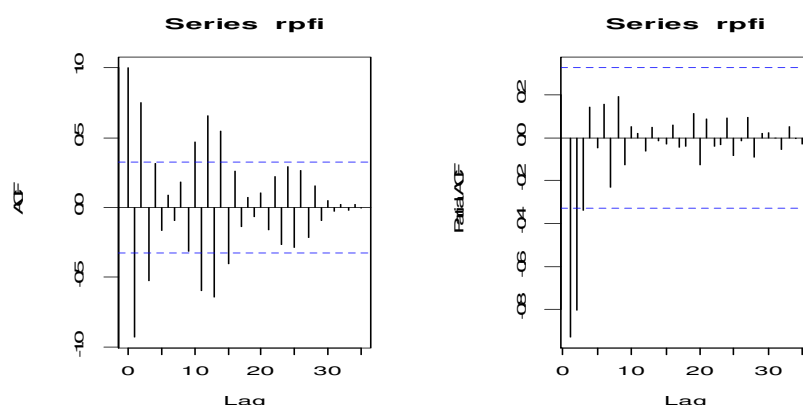
Figura 8 – Série PFI com diferença 12



Fonte - O autor.

Um novo teste de raiz unitária para essa série com diferença 12 foi realizado e, desta vez, foi rejeitada a hipótese nula de não estacionariedade, com p-valor de 0,01.

Agora, com uma série estacionária, daremos início ao processo de estimação de um modelo SARIMA para o conjunto de dados. Os gráficos de fac e de facp da série PFI com diferença 12 seguem abaixo, na Figura 9.

Figura 9 – FAC e FACP da Série PFI com diferença 12

Fonte - O autor.

Inspecionando os gráficos fac e facp da série com diferença 12, é possível perceber que os primeiros lags (1,2,...) são estatisticamente significantes. Desse modo, de acordo com o final da seção 3 e baseado em Shumway and Stoffer (2011), podemos sugerir $p = 1$ e $q = 1$ ou $p = 3$ e $q = 0$.

Por outro lado, a facp decresce exponencialmente enquanto que a fac apresenta lags significantes até na 15ª posição, sendo que o 12º lag é o que possui (em módulo) o maior valor. Nesse caso, temos que $k = 1$ e $s = 12$.

Dessa forma, por considerar o gráfico da facp como decrescente, enquanto que a fac trunca no lag 12 (1 período de sazonalidade), vamos sugerir $Q = 1$ e $P = 0$ e o modelo SARIMA poderá admitir uma das formas a seguir:

- $ARIMA(1,0,1) \times (0,1,1)_{12}$;
- $ARIMA(3,0,0) \times (0,1,1)_{12}$

Quando da estimativa de ambos os modelos, após os ajustes necessários (remoção de variáveis estatisticamente não significantes), ambos os modelos convergiram para um mesmo modelo em comum, a saber, o modelo $ARIMA(1,0,0) \times (0,1,0)_{12}$, com coeficiente $\phi_1 = 0,9258$ (e.p. 0,0589).

Com relação ao teste diagnóstico, os resíduos do modelo estimado $ARIMA(1,0,0) \times (0,1,0)_{12}$ foram testados com o objetivo de verificar a existência de auto-correlação. O resultado do teste Box-Ljung sugere a não rejeição da hipótese nula, com p-valor de 0,8076, assegurando a adequação do modelo estimado.

Uma vez encontrado um modelo SARIMA adequado para o conjunto de dados, podemos dar início ao processo de análise de intervenção. A função de transferência mais apropriada aos nossos interesses é dada pela expressão (4.4), ou seja,

$$v(B)X_t = \begin{cases} 0, & t \neq T \\ \frac{\omega_2 B}{1 - \delta B}, & t = T. \end{cases}$$

A escolha por ω_2 deve-se ao fato de que apesar da escolha para o ponto do tempo do evento de interesse tenha sido especificamente o mês de setembro, o primeiro movimento de queda na série PFI ocorreu somente no segundo mês subsequente.

Seguem, portanto, os resultados dessa modelagem na Tabela 2.

Tabela 2 – Estimação dos Parâmetros da Análise de Intervenção

Coeficientes	ar1	T1-AR1	T1-MA0	T1-MA1	T1-MA2
Estimativa	0,5793	0,8875	-1,3167	-7,6900	-14,4325
S.E.	0,2097	0,0234	3,9867	4,1493	5,2416
sigma^2 estimated as 15.75: log likelihood = -95.31, aic = 200.62					

Fonte: O autor.

Com relação aos testes de significância para os parâmetros estimados na análise de intervenção, é possível afirmar, ao nível de $\alpha = 0,05$, que:

- O parâmetro δ é estatisticamente significativo ($t_{cal} = \frac{0,8875}{0,0234} = 37,927$);
- O parâmetro ω_2 é estatisticamente significativo ($t_{cal} = \frac{-14,4325}{5,2416} = -2,75$);
- O parâmetro ω_1 não se mostrou estatisticamente significativo ($t_{cal} = \frac{-7,69}{4,1493} = -1,85$);
- O parâmetro ω_0 não se mostrou estatisticamente significativo ($t_{cal} = \frac{-1,3167}{3,9867} = -0,33$).

Com base nos resultados acima, é possível notar que apenas os parâmetros δ e ω_2 se revelaram estatisticamente significantes, ao nível de 5%, permitindo que a função de transferência de interesse possa ser escrita como

$$v(B)X_t = \begin{cases} 0, & t \neq T \\ \frac{-14,4325}{1 - 0,8875B}, & t = T. \end{cases}$$

Ainda sobre os resultados dos testes acima, alguns comentários sobre ω_0 e ω_1 merecem destaque. Os resultados de seus testes de significância corroboraram a suspeita inicialmente observada no gráfico da Figura 10, a

saber, que a crise não teria produzido efeitos na PFI, seja imediatamente (ω_0), seja no período subsequente (ω_1).

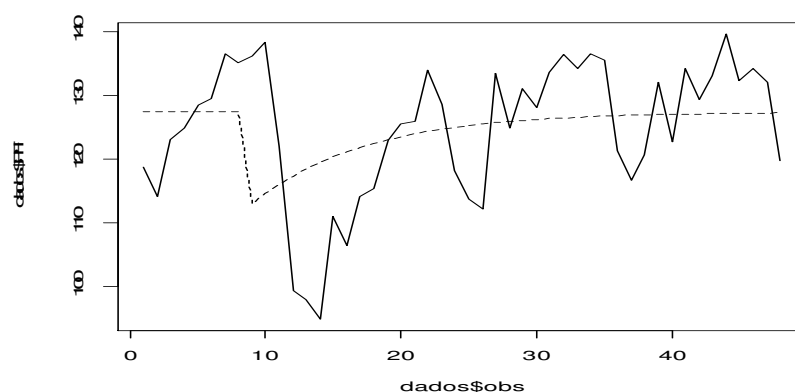
É possível notar, também, que o sinal do parâmetro estimado para ω_2 é negativo, o que admite uma interpretação mais intuitiva, uma vez que era esperado que a intervenção (crise) influenciasse negativamente a trajetória da curva.

Wei (2006), por outro lado, especifica que quando $\delta = 1$, o impacto é linear e ilimitado, mas, que, na maioria dos casos, $0 < \delta < 1$, refletindo uma resposta mais gradual. Uma vez, pelo modelo ajustado, que $\hat{\delta} = 0,8875$, pode-se pressupor que a recuperação da economia brasileira tenha ocorrido de forma gradual, mas, com estimativa não muito distante da tendência linear, sugerindo, desse modo, que a recuperação da atividade econômica ocorreu de forma proeminente.

Box, Jenkins e Reinsel (1994) sugerem uma interpretação mais geral, tal que o impacto após a intervenção em setembro, com defasagem de 2 meses, tem ocorrido em uma magnitude -14,4325 pontos (uma vez que se trata de uma série temporal do tipo “Número Índice”), seguida por um crescimento gradual com taxa 88,75% de volta para o nível original pré-intervenção, sem efeitos permanentes. A visão gráfica dos efeitos da intervenção na série de PFI e a subsequente recuperação da série PFI podem ser observadas no gráfico da Figura 10.

Por fim, o teste Box-Ljung referente aos resíduos do modelo contido na Tabela 2 produziu p-valor 0,6763, sugerindo inexistência de autocorrelação residual e, portanto, adequação do modelo ao conjunto de dados.

Figura 10 – Efeitos da Intervenção na Série PFI



Fonte: O autor.

6. CONCLUSÕES

Os resultados da análise de intervenção reforçam a hipótese de que a crise bancária americana produziu, de forma significativa, efeitos negativos na produção industrial brasileira.

Na análise de intervenção foi escolhida a função de intervenção que representasse a característica de defasagem entre a crise e sua efetiva captação nos índices de produção nacional. Os resultados obtidos mostraram consistência, uma vez que as estimativas ω_0 , ω_1 não se mostraram estatisticamente significantes, enquanto ω_2 , por outro lado, mostrou-se estatisticamente significativo (todas ao nível de 0,05). Em outras palavras, há evidências de que os efeitos da crise não foram transmitidos de forma simultânea/automática para a economia brasileira. Houve uma defasagem entre o período da deflagração da crise (setembro de 2008) e os impactos iniciais percebidos na própria série temporal de produção física industrial - PFI, sendo que essa defasagem, a qual foi incorporada ao modelo, revelou-se estatisticamente significativa. Além disso, há evidências de que a recuperação da atividade econômica (logo após o impacto negativo da crise) deu-se de forma proeminente, ainda que de forma gradualista ($\hat{\delta} = 0,8875$), relativamente próxima de uma tendência linear.

7. BIBLIOGRAFIA

BOX, G. E. P., JENKINS, G. M. and REINSEL, G. C. (1994). "Time Series Analysis: Forecasting and Control". 3rd edition. Prentice-Hall International, Inc. 1994.

BOX, G. E. P. and TIAO, G. C. (1975). "Intervention Analysis with Applications to Economic and Environmental Problems". Journal of American Statistical Association, Vol. 70, No. 349 (Mar., 1975), 70-79.

BRESSER-PEREIRA, L. C. (2009). "A Crise Financeira de 2009". Revista de Economia Política, vol. 29, nº 1 (113), pp. 133-149, janeiro-março/2009.

CECHIN, A. e MONTTOYA, M. A. (2017) Origem, causas e Impactos da crise financeira de 2008. Teoria e Evidência Econômica – Ano 23, n. 48, p.150-171, jan/jun.

CRYER, J. D. e CHAN, K. S. (2008). "Time Series Analysis with Applications in R". Springer Texts in Statistics. 2nd Edition. 2008.

FAVA, V. L. e ALVES, D. C. O. (1997). "Indicador de Movimentação Econômica, Plano Real e Análise de Intervenção". RBE. Rio de Janeiro. 51(1):133-43. Jan/Mar. 1997.

FREITAS, M. C. P. de (2009) Os efeitos da crise global no Brasil: aversão ao risco e preferência pela liquidez no mercado de crédito. Estud. Av. vol.23 no.66 São Paulo.

GUJARATI, D. (2006). "Econometria Básica". Rio de Janeiro. Elsevier. Campus. 2006.

JUNIOR, G. R. B. e FILHO, E. T. R. (2008). "Analisando a Crise do Subprime". Revista do BNDES, Rio de Janeiro, V. 15, N. 30, P. 129-159, Dez.2008.

KRUGMAN, P. (2009). "The Return of Depression Economics and the Crisis of 2008". W. W. Norton & Company. New York. London. 2008.

LEMOES, B. P. e BITTENCOURT, M. V. L. (2007). "A Crise Imobiliária Americana: Origem, Perspectivas e Impactos". Economia & Tecnologia. Ano 03, Vol. 10 – Jul./Set. de 2007.

LJUNG, G. M. and BOX, G. E. P. (1978). "On a Measure of Lack of Fit in Time Series Models". Biometrika, 65, 2, pp. 297-303.

MORETTIN, P. A. (2006). "Econometria Financeira: Um Curso em Séries Temporais Financeiras". 17º SINAPE - ABE. 2006.

MORETTIN, P. A. e TOLOI, C. M. C. (2006). "Análise de Séries Temporais". 2ª ed. – São Paulo: Edgard Blucher. 2006.

NEUSSER, K. (2016) Time Series Econometrics. Springer International Publishing Switzerland. 2016.

R Core Team (2019). R: A language and environment for statistical computing. R

Foundation for Statistical Computing, Vienna, Austria. URL

<https://www.R-project.org/>.

SANTIAGO, M. M. D., CAMARGO, M. de L. B. e MARGARIDO, M. A. (1996). "Detecção e Análise de Outliers em Séries Temporais de Índices de Preços Agrícolas no Estado de São Paulo". Agricultura em São Paulo, SP, 43(2):89-115, 1996.

SHUMWAY, R. and STOFFER, D. S. (2011). "Time Series Analysis and Its Applications With R Examples". 3rd Edition, Springer Texts in Statistics. 2011.

TERAZI, E. and SENEL, S. (2011) "The effects of the Global Financial Crisis on the Central and Eastern European Union Countries". International Journal of Business and Social Science. Vol. 2, No. 17.

TUPY, I. S. (2015) Impactos Regionais de Crises Financeiras: Estudo sobre as respostas dos Estados Brasileiros à Crise Financeira Global. Dissertação apresentada ao curso de Mestrado em Economia do Centro de Desenvolvimento e Planejamento Regional Da Faculdade de Ciências Econômicas da Universidade Federal de Minas Gerais.

WEY, W. W. S. (2006) Time series analysis: univariate and multivariate methods. 2nd ed. Pearson Education, Inc. 2006.

SITES:

S&P/Case-Shiller (2012). "Home Price Indices: 2011 Year in Review". Disponível em:

<<http://us.spindices.com/index-family/real-estate/sp-case-shiller>>

www.ibge.gov.br

<https://www.ipeadata.gov.br>

USO DE FERRAMENTAS ECONOMETRICAS PARA MODELAR E ESTIMAR O PIB DO BRASIL

Waslley Pablo Costa da Silva

waslleypablo@hotmail.com

Universidade Estadual de Campinas

Fernando Ferraz do Nascimento

fernandofn@ufpi.edu.br

Universidade Federal do Piauí

Valmária Rocha da Silva Ferraz

valmaria@ufpi.edu.br

Universidade Federal do Piauí

Resumo: O Brasil ainda é considerado um país de renda média, apesar do crescimento econômico e social dos últimos anos, estima-se que apenas em 2035 seremos minimamente desenvolvidos, e disso deve ser reduzido em vista de nossas potencialidades e riquezas. Entenda como as variáveis econômicas ou sociais que influenciam a compreensão do PIB são uma das alternativas para encontrar modelos econométricos que possam ajudar no planejamento da tomada de decisões de natureza econômica e, assim, construir o caminho para um país desenvolvido. Este trabalho consiste em estimar o PIB do Brasil para o ano de 2019, usando uma comparação de vários modelos de séries temporais. Os resultados dos modelos mostram que o método foi eficiente nas projeções de 2018, comparado às projeções oficiais e extra oficiais. Além disso, a previsão para 2019 também se aproxima dessas projeções oficiais.

Palavras-chave: Séries Temporais; Econometria; Box-Jenkins; ARIMAX; SARIMAX; PIB

Abstract: Brazil is still considered a middle-income country despite economic and social growth of the last years, it is estimated that only in 2035 will we be minimally developed, this needs to be shortened in view of our potentialities and riches. Understand how economic or social variables that influence the understanding of GDP is one of the in order to find econometric models that can help in planning decision-making of an economic nature and thus build the path to be a developed country. This work consists of estimating Brazil's GDP for the year 2019 using a comparison of various time series models. Results of the models showed that the method was efficient in the projections of 2018, compared to official and extra official projections. In addition, the 2019 forecast is also close to these official projections.

Keywords: Time Series; Econometry; Box-Jenkins; ARIMAX; SARIMAX; GDP

1.INTRODUÇÃO

O PIB é a soma de todos os bens e serviços finais da economia produzidos num determinado país, medido em unidades monetárias. É um indicador importante na tomada de decisões econômicas e na formulação de políticas públicas, além de mostrar o crescimento ou decréscimo da produção interna de uma nação, entes federativos e/ou municípios.

No Brasil o PIB é calculado e divulgado oficialmente pelo IBGE (Instituto Brasileiro de Geografia e Estatística), órgão vinculado ao ministério do planejamento e ao governo federal do Brasil criado em 1934 e instalado em 1936. A metodologia do IBGE para o cálculo do produto advém de uma cartilha das nações unidas, denominada *System of National Accounts*. A divulgação do produto é dada trimestralmente e anualmente, no caso trimestral há um atraso de três meses, ou seja, em cada trimestre se divulga o PIB do trimestre imediatamente anterior, o IBGE usa o primeiro dia do último mês de cada trimestre como data oficial para a divulgação desse indicador.

Para analisarmos os dados que serão apresentados e posteriormente modelarmos os mesmos se fará necessário o uso de uma ferramenta chamada econometria. Segundo Gujarati & Porter (2011, p. 25) "a econometria pode ser definida como a análise quantitativa dos fenômenos econômicos ocorridos com base no desenvolvimento paralelo da teoria e das observações e com o uso de métodos de inferência adequados".

O Brasil ainda é um país considerado de renda média, apesar do avanço social e econômico dos últimos anos. O PIB *per capita* do Brasil em 2017 chegou a R\$31.587,00. Esse valor ainda é muito pouco se comparado ao PIB per capita de nossos parceiros comerciais desenvolvidos, estima-se que somente em 2035 nós brasileiros iremos ter um padrão de vida razoavelmente elevado, se crescermos no mínimo 2,8% ao ano, pelos próximos anos, segundo o documento "Visão 2035: Brasil, país desenvolvido." do BNDES (2018).

A depressão econômica 2014-2016 nos deixou ainda mais distante de sermos desenvolvidos, precisamos de modelos econométricos, que nos permita ter uma base da importância das variáveis macroeconômicas para a estimação do PIB, que permita uma tomada de decisão correta afim do desenvolvimento e do crescimento econômico.

1.1 OBJETIVOS

O objetivo consiste em descobrir qual modelo econométrico se ajusta aos dados anuais e trimestrais respectivamente, e posteriormente usar o modelo escolhido para previsão. Serão aplicados alguns modelos de séries temporais, ARIMA e SARIMA, com a presença de co-variáveis.

2. METODOLOGIA BOX-JENKINS

Para analisar modelos paramétricos, utiliza-se a metodologia Box & Jenkins (1970), conhecido como ARIMA(p,d,q). Esse método se utiliza das propriedades probabilísticas e estocásticas das séries temporais. Nos modelos ARIMA, Y_t é explicado pelos termos anteriores $Y_{t-1}, \dots, Y_{t-k}$ e pelo erro aleatório. O passo para a construção dos modelos se dá de forma iterativa, consistido em 4 etapas, primeiramente se escolhe uma classe geral de modelos para verificação, em seguida define-se um modelo com base nas autocorrelações amostrais e parciais e é escolhido seus parâmetros, depois vem a análise residual, se o modelo for aprovado nessa etapa ele vai para a última etapa se não volta-se para a fase de identificação do modelo, por fim temos a previsão do modelo proposto.

2.1 MODELOS SARIMA

Segundo Morettin & Tolo (2006) se observa na prática que as séries não são estacionárias, um exemplo clássico são as séries temporais econômicas, que geralmente são tendenciosas. De acordo com Bezerra (2006), para tornar estacionária uma série temporal recomenda-se aplicar a diferença quantas vezes for necessário, até alcançar a estacionariedade.

O procedimento para a realização das diferenças é dado a seguir:

$$\Delta_t = Y_t - Y_{t-1} = (1 - B)Y_t \quad (1)$$

Dizemos Y_t é uma série não-estacionária homogênea quando, podemos aplicar um número finito de diferença d. Logo:

$$W_t = \Delta^d Y_t \quad (2)$$

Segundo Gerolimetto (2010), as séries temporais podem apresentar uma componente repetida a cada s intervalos de tempo, devido essa repetição a cada instante s. Assim, o modelo, chamado de SARIMA(p,d,q)(P,D,Q)[s], pode ser descrito da seguinte maneira:

$$\Phi_p(B^S)\Phi(B)\Delta_S^D\Delta^d Y_t = \theta_Q(B^S)\theta(B)\alpha_t \quad (3)$$

Onde:

$\Phi(B)$ é o operador autorregressivo de ordem p, estacionário;

Δ^d é a diferença de ordem d ;

$\theta(B)$ é o operador de médias móveis de ordem q , irreduzível;

$\Phi_P(B^S)$ é o operador autorregressivo sazonal de ordem P ;

Δ_S^D é a diferença sazonal de ordem D ;

$\Theta_Q(B^S)$ é o operador de médias móveis sazonal de ordem Q ;

α_t é um processo de ruído branco;

Y_t é a série temporal principal.

2.1 MODELOS SARIMAX

Segundo Medeiros (2018), os modelos SARIMA(p,d,q)x(P,D,Q)[S] se utilizam de seus valores passados afim de fazer previsões para valores futuros, contudo em certos casos essas previsões podem não ser robustas pois algumas séries temporais dependem de determinados fatores contidos em outras séries temporais. Nesse sentido de dar mais acuracidade a previsão e na própria modelagem da série, surgiu a necessidade de acrescentar aos modelos SARIMA variáveis exógenas, que se intitulou SARIMAX(p,d,q,x)x(P,D,Q,X)[S] para séries sazonais, essa modelagem foi instituída por Box & Tiao (1975), suas fórmulas são dadas a seguir respectivamente:

$$\Phi_P(B^S)\Phi(B)\Delta_S^D\Delta^d(Y_t - x_t\beta) = \theta_Q(B^S)\theta(B)\alpha_t \quad (4)$$

Onde:

x_t são as séries temporais exógenas;

β são os coeficientes das séries temporais exógenas.

3. APLICAÇÃO

Para análise dos bancos de dados e para a construção dos modelos econométricos, foi utilizado o *software* R (R Core Team, 2018), e os seguintes pacotes: *lmtest* (Horthorn *et al.*, 2019), *forecast* (Hynkman & Khandakar, 2008), *urca* (Pfaff *et al.*, 2016), *tseries* (Trapletti *et al.*, 2019) e *TSA* (Chan & Ripley, 2018). Foi utilizado os seguintes testes para adequação dos modelos as suposições dos métodos que serão utilizados: Philip-Perron (1988) para presença de raiz unitária e Ljung-Box-Pierce (1978) para a independência dos resíduos, afim de verificar os pressupostos do método Box & Jenkins.

Os dados são originários do IPEADATA e do BCB, estão disponíveis nos sítios <http://www.ipeadata.gov.br> e <https://dadosabertos.bcb.gov.br>, respectivamente. O estudo é do tipo longitudinal e foi utilizado para analisar os dados nos métodos de séries temporais utilizando-se da metodologia Box &

Jenkins (1970), no período de 1996 a 2017 e no período de 1996 a 2018, com periodicidade anual e trimestral, respectivamente. Primeiro foram obtidas as estatísticas descritivas de cada variável, posteriormente a construção dos modelos e escolha dos mesmos. As variáveis são oriundas do IPEADATA e BCB, no período de 1996 a 2017 e no período de 1996 a 2018, a seguir exemplificamos as variáveis que serão aqui analisadas.

- PIB: É a soma dos bens e serviços de um país produzidos em um certo período de tempo, mostra a realidade da economia em determinado período, o PIB é divulgado trimestralmente pelo IBGE pelo sistema de contas nacionais;
- Consumo das Famílias: De acordo com Gomes (2012) é o consumo de bens e serviços finais da economia afim de satisfazer as necessidades dos indivíduos, é calculado e divulgado pelo IBGE trimestralmente pelo sistema de contas nacionais;
- ICC – Índice de Confiança do Consumidor Paulistano: Índice que mostra os sentimentos sobre a situação econômica atual, varia de 0 a 200, sendo abaixo de 100 pessimismo e acima de 100 otimismo é divulgado e calculado pela Fecomércio de São Paulo desde o começo da década de 1990, sua metodologia se baseia indicador de confiança de Michigan criado em 1950;

3.1 ESTATÍSTICAS DESCRITIVAS

Tabela 1 - Estatísticas Descritivas dos Dados Anuais e trimestrais excluso 2018

Dados Anuais							
Variável	Mín	1º Quartil	Mediana	Média	3º Quartil	Máx	CV (%)
PIB	3277655	3858676	4956373	5202492	6765309	7303581	28.22
Consumo privado	2133123	2448810	2970675	3236025	4213293	4661379	27.99
ICC ¹	90.40	99.05	113.15	120.04	138.88	161.80	19.6
Dados Trimestrais							
Variável	Mín	1º Quartil	Mediana	Média	3º Quartil	Máx	CV (%)
PIB	755239	951590	1248140	1300623	1692381	1884286	27.98
Consumo privado	500958	611084	745881	809006	1048714	1211334	27.68
ICC	81.6	101.1	112.7	120.0	136.5	164.3	20.1

Fonte: O Autor

Tabela 2 - Estatísticas Descritivas dos Dados Anuais e trimestrais incluso 2018

Dados Anuais							
Variável	Mín	1º Quartil	Mediana	Média	3º Quartil	Máx	CV (%)
PIB	3323722	3954206	5229242	5354958	6915469	7518635	28.15
Consumo privado	2166082	2490960	3131281	3340972	4380727	4733402	28.15
ICC ¹	90.4	99.6	111.8	119.04	137.1	161.80	19.26
Dados Trimestrais							
Variável	Mín	1º Quartil	Mediana	Média	3º Quartil	Máx	CV (%)
PIB	766269	967866	1310643	1338740	1720505	1924767	27.92
Consumo privado	508699	623220	780247	835243	1091151	1230050	27.73
ICC	81.6	101.9	112.7	119.7	136.1	164.3	19.78

Fonte: O Autor

Podemos observar nas Tabelas 1 e 2 e nas Figuras 1 e 4, temos que o mínimo valor do PIB, se refere ao início da série do mesmo em 1996 e o máximo em 2014, assim como o consumo privado, já a confiança do consumidor tem o seu valor mínimo em 1999 quando chega a 90.4 e um máximo de 161.80 no ano de 2012. Empiricamente podemos notar que a confiança do consumidor tende a ser mais sensível aos impactos da política e da economia, como acontece por exemplo no ano de 1999 com a desvalorização do câmbio, nos anos de 2015-2016 e em anos eleitorais. Contudo para entendermos melhor a dinâmica de comportamento dessas variáveis e como o PIB se comporta em função das outras, precisamos criar modelos que nos permitam um diagnóstico confiável.

No caso trimestral conseguimos ver com mais detalhes nas Tabelas 1 e 2 e nas Figuras 1 e 4 a partir de qual período de tempo houve uma mudança brusca no comportamento das séries, o valor mais alto para o PIB trimestral é do último trimestre de 2014, assim como o consumo familiar, já a confiança do consumidor atingiu seu valor máximo no primeiro trimestre de 2012 e mínimo no segundo trimestre de 1999. Assim como nos dados anuais o ICC tende a ser mais sensível aos choques políticos e econômicos.

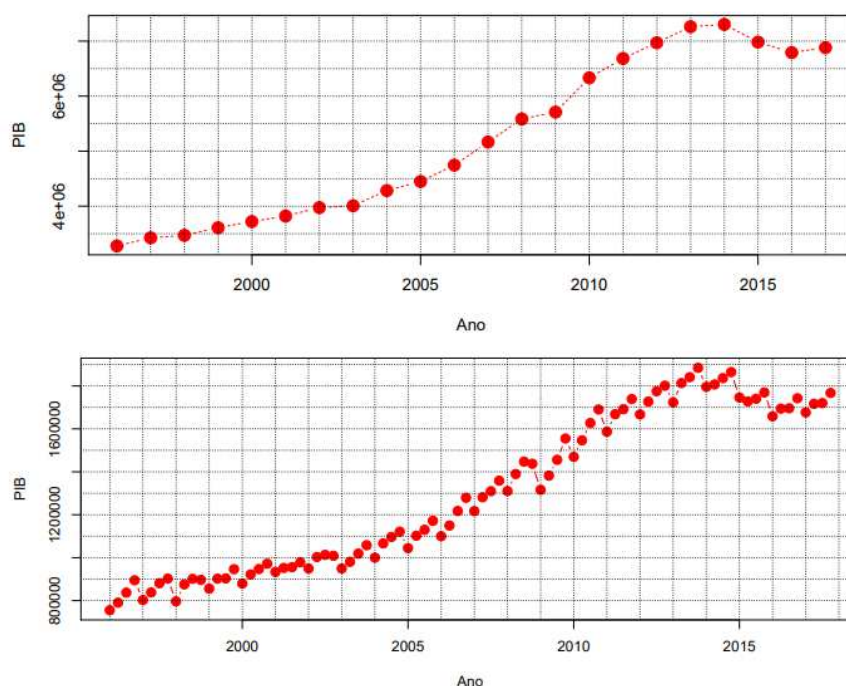
3.2 SÉRIES TEMPORAIS DAS VARIÁVEIS

Outro recurso bastante útil para a análise econométrica são os gráficos das variáveis ao longo do tempo, eles nos permitem uma visão empírica do comportamento da variável em cada instante de tempo. Nas Figuras 1, 2 e 3 temos as séries das variáveis aqui estudadas no período de 1996-2017 e nas Figuras 4, 5 e 6 no período 1996-2018, o PIB e o Consumo Privado foram

deflacionadas pelo IPCA do corrente ano em que se termina cada série e calculado pelo IPEADATA, podemos notar que as séries do PIB e do consumo privado, praticamente possuem o mesmo comportamento, esse fenômeno ocorre porque o PIB é composto pelo consumo privado, ou seja qualquer alteração no consumo privado irá refletir no PIB.

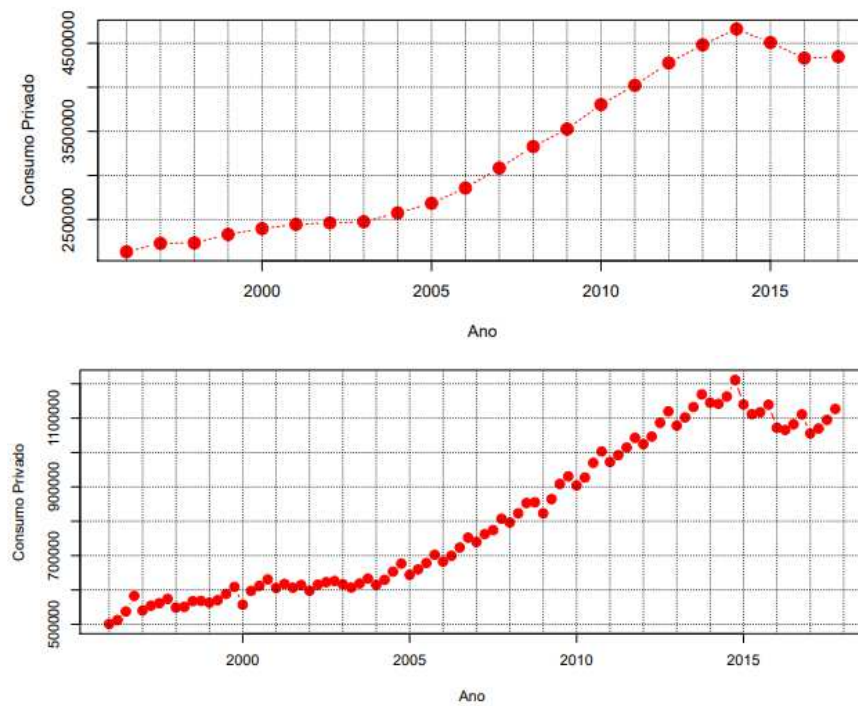
Ao olharmos as Figuras 3 e 6, percebemos que o ICC tendeu a variar mais bruscamente que as outras variáveis, podemos visualizar que o ICC atingiu o máximo histórico em 2012 e a partir daí começou a cair, e de forma brusca entre 2013 a 2015, período de grande instabilidade política e econômica. Segundo Aranha (2017), a recessão brasileira trouxe uma nova perspectiva para os indicadores de confiança, como o ICC. Em meados de 2013, quando o crescimento econômico ainda não apresentava nítida desaceleração, os indicadores de confiança do consumidor, tais como o ICC, já indicavam relevante piora. Em 2014, quando a atividade estava desacelerando, estes indicadores continuaram caindo até tocar seu ponto mínimo em 2015. Contudo, em 2016, observou-se uma melhora destes indicadores de confiança, tais como o ICC paulistano. Podemos notar também passada a instabilidade política econômica causada pelo *impeachment* da ex-presidente Dilma Rousseff, o ICC começou a se recuperar e logo depois o PIB e o Consumo também começaram a se recuperar como mostraram as Figuras 1, 2, 4 e 5.

Figura 1: PIB Anual 1996-2017 / PIB Trimestral 1996.1-2017.4



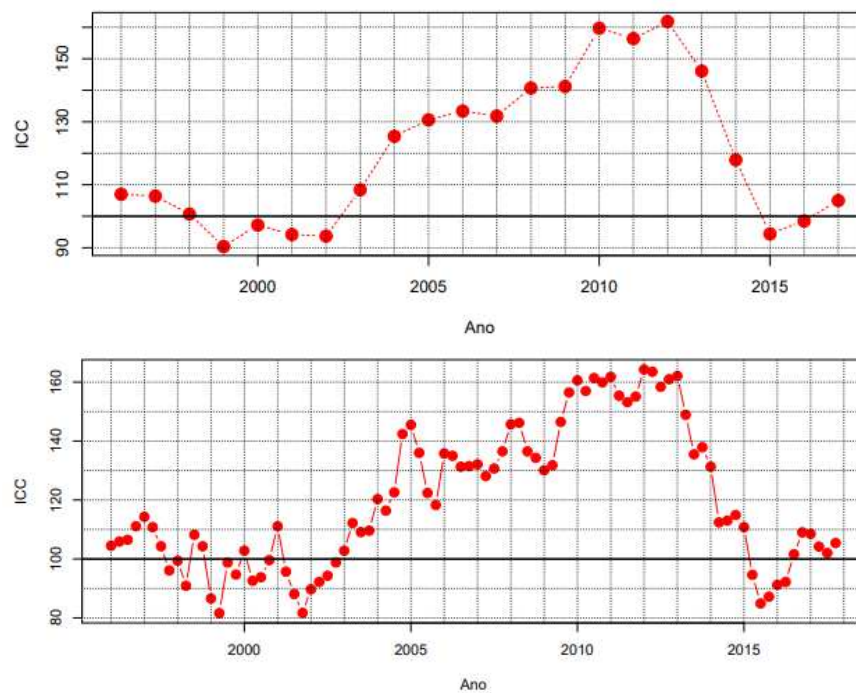
Fonte: O Autor

Figura 2: Consumo Privado Anual 1996-2017 / Consumo Privado Trimestral 1996.1-2017.4



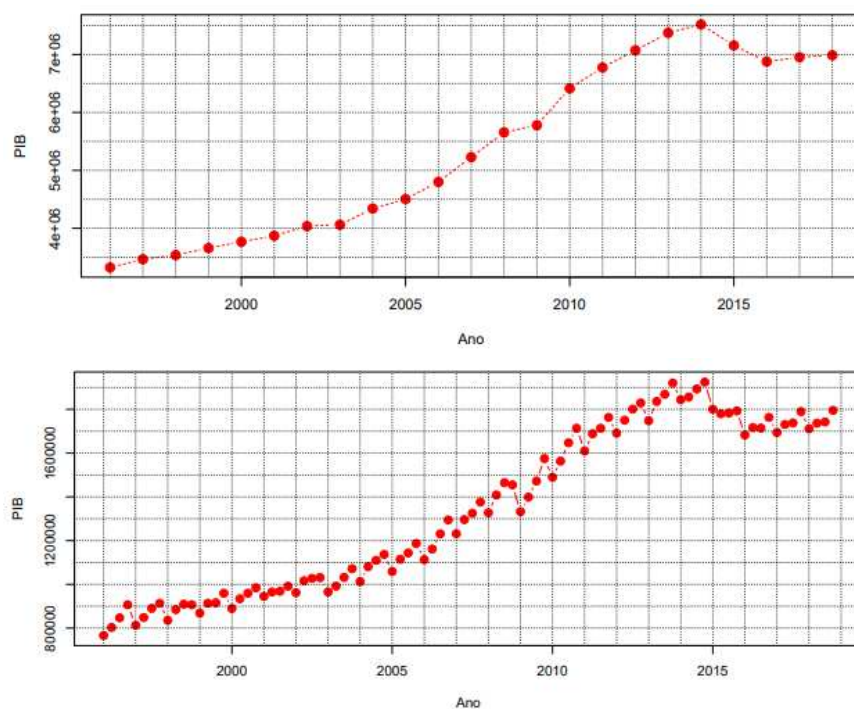
Fonte: O Autor

Figura 3: ICC Anual 1996-2017 / ICC Trimestral 1996.1-2017.4



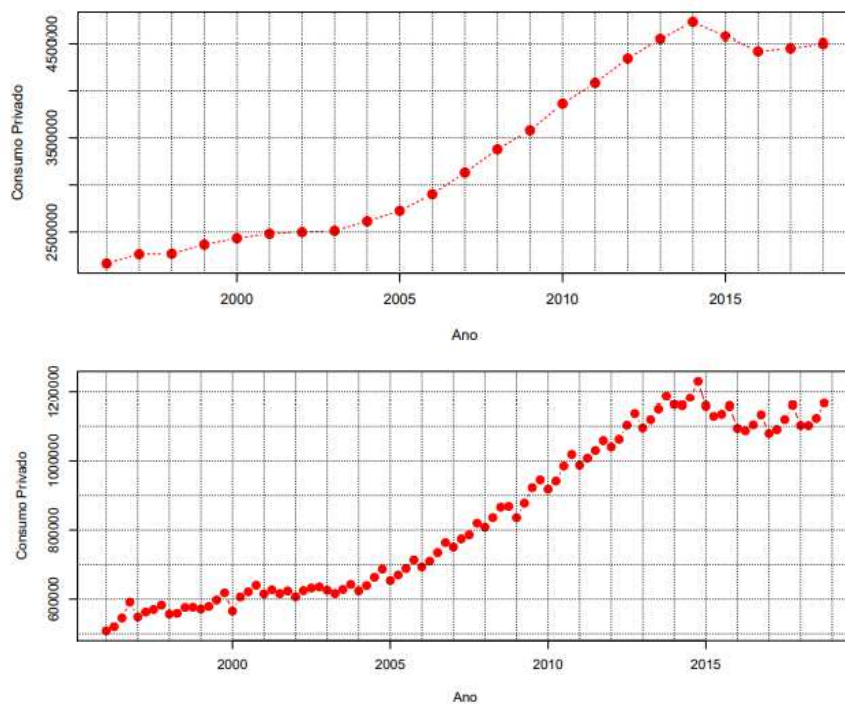
Fonte: O Autor

Figura 4: PIB Anual 1996-2018 / PIB Trimestral 1996.1-2018.4



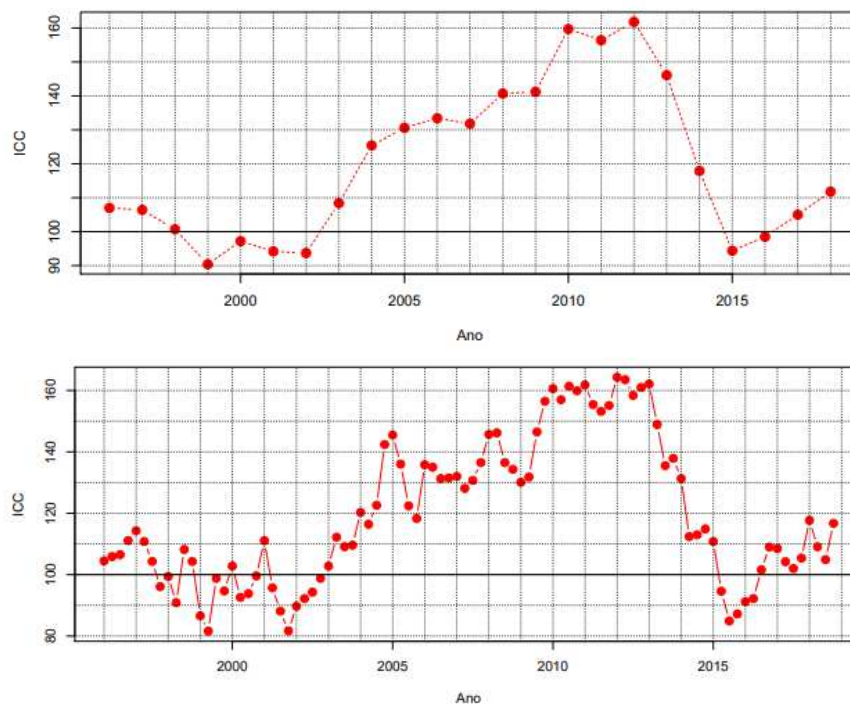
Fonte: O Autor

Figura 5: Consumo Privado Anual 1996-2018 / Consumo Privado Trimestral 1996.1-2018.4



Fonte: O Autor

Figura 6: ICC Anual 1996-2018 / ICC Trimestral 1996.1-2018.4

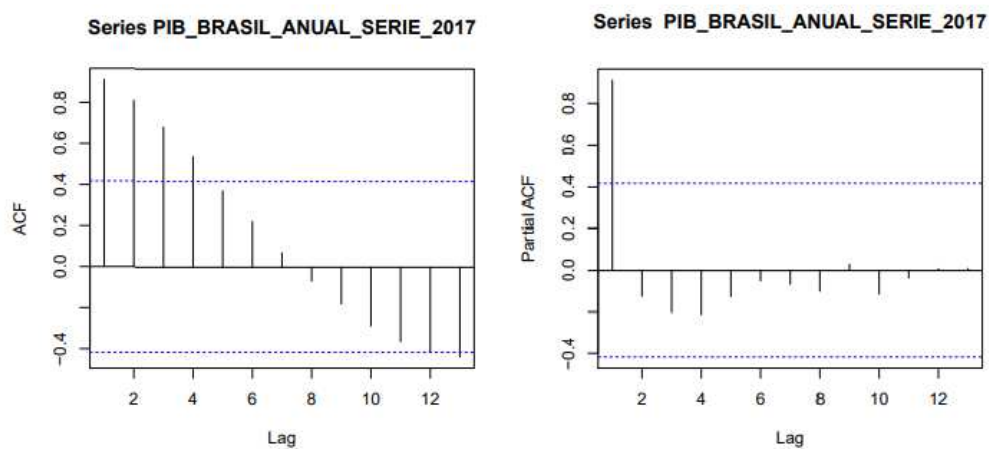


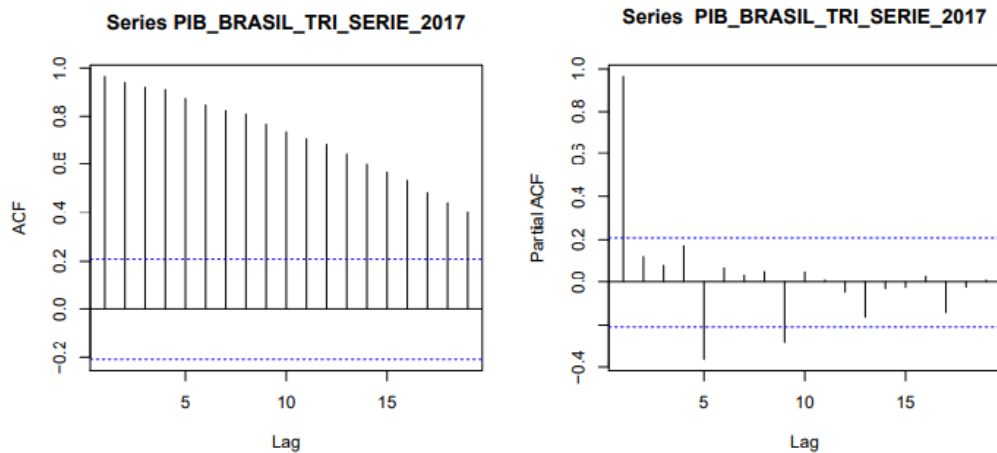
Fonte: O Autor

3.3 FAC E FACP DAS VARIÁVEIS

Construímos os gráficos da FAC e FACP para o PIB, Consumo familiar privado e ICC, anual e trimestral respectivamente, como mostram as Figuras 7, 8, 9, 10, 11 e 12, todas as séries se mostraram não estacionárias necessitando de duas diferenças para as séries anuais e uma diferença para as trimestrais.

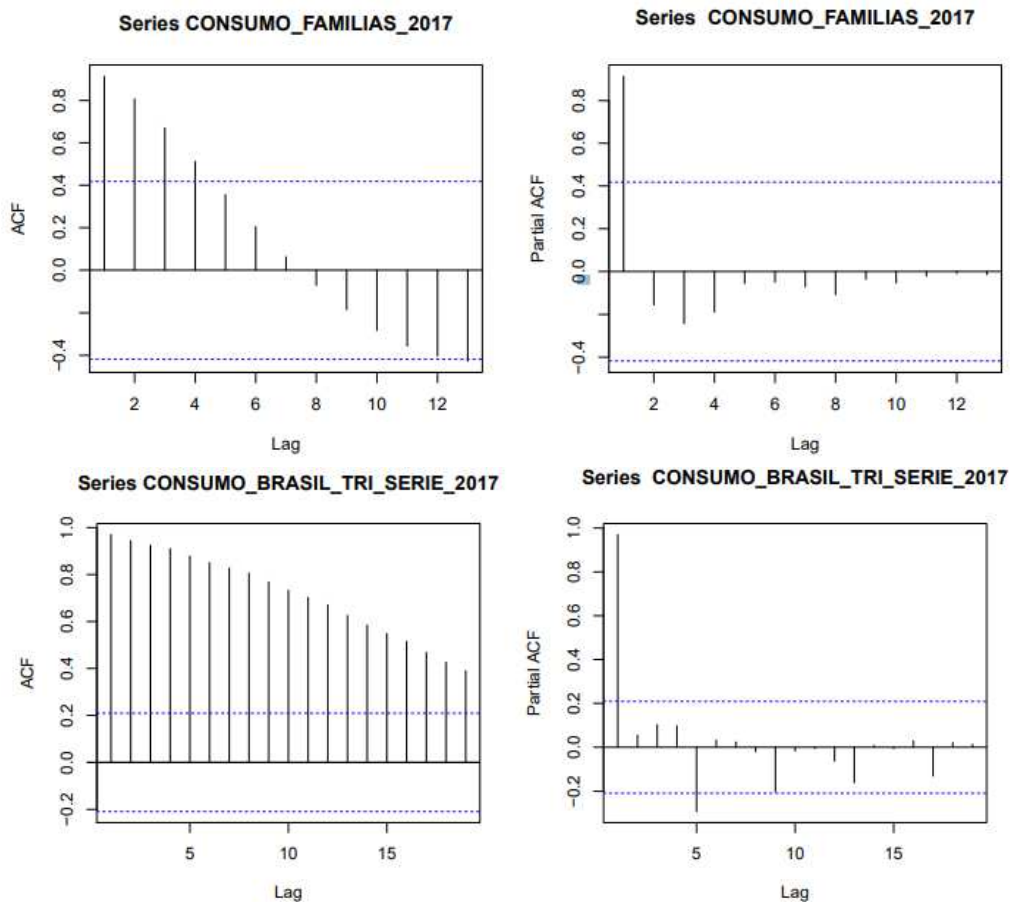
Figura 7: FAC e FACP do PIB anual de 1996 a 2017 / FAC e FACP do PIB trimestral de 1996.1 a 2017.4





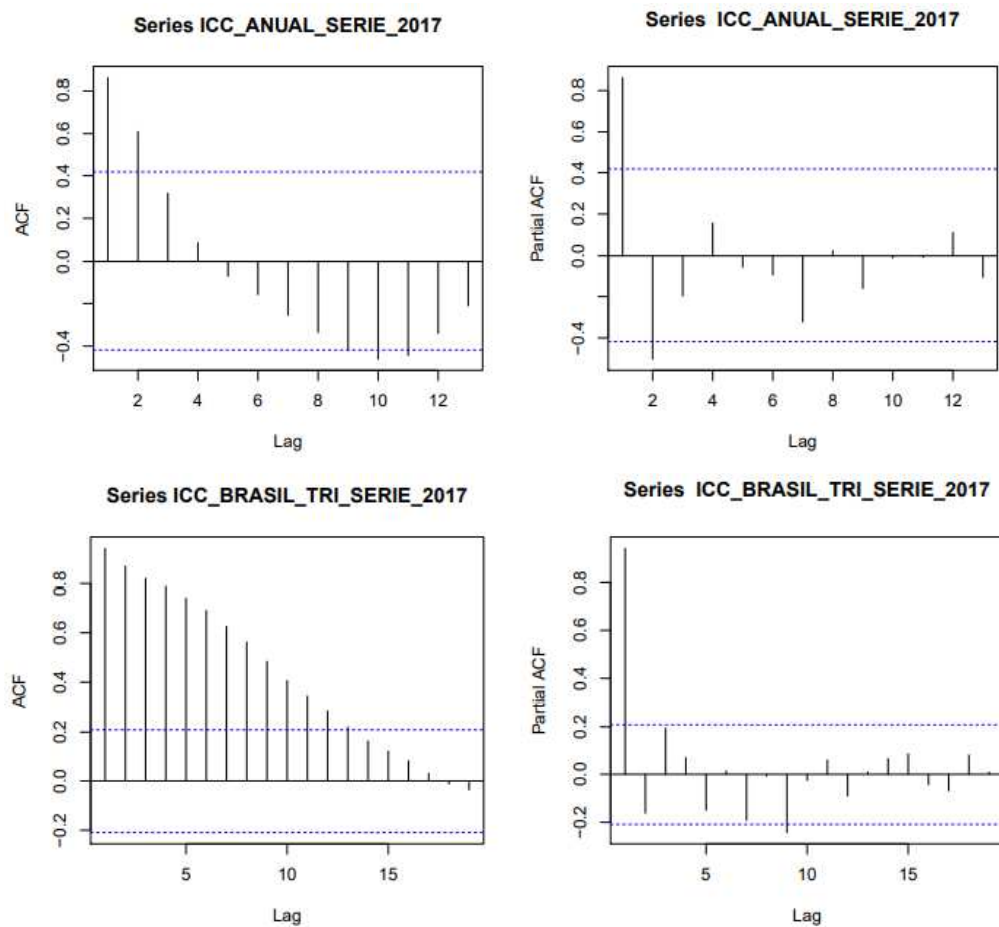
Fonte: O Autor

Figura 8: FAC e FACP do Consumo Privado anual de 1996 a 2017 / FAC e FACP do ICC
Consumo Privado trimestral de 1996.1 a 2017.4



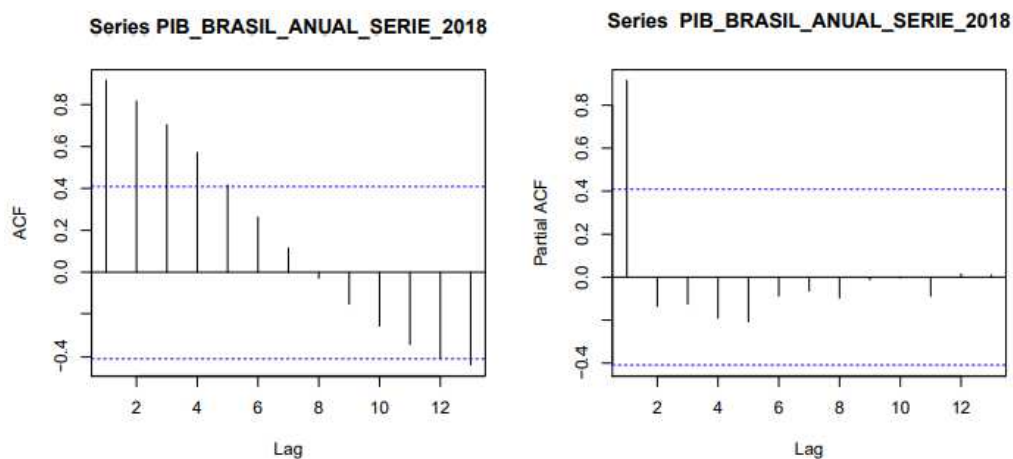
Fonte: O Autor

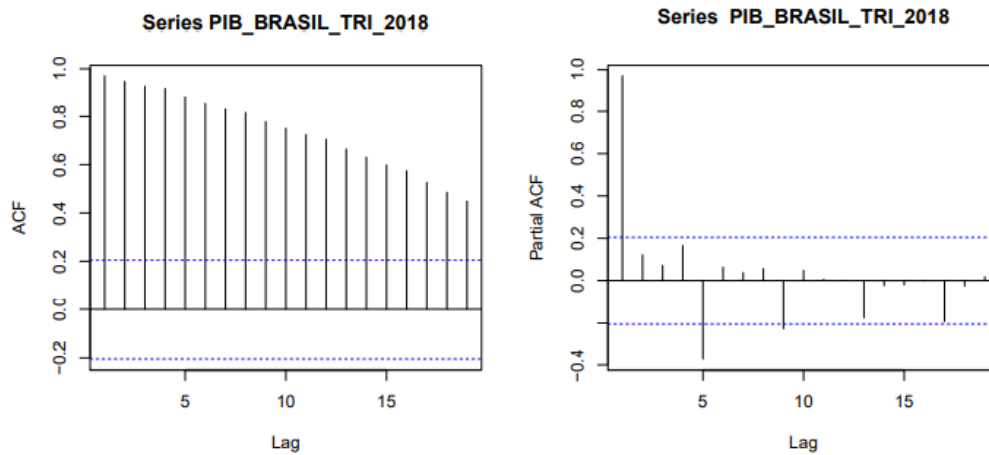
Figura 9: FAC e FACP do ICC anual de 1996 a 2017 / FAC e FACP do ICC trimestral de 1996.1 a 2017.4



Fonte: O Autor

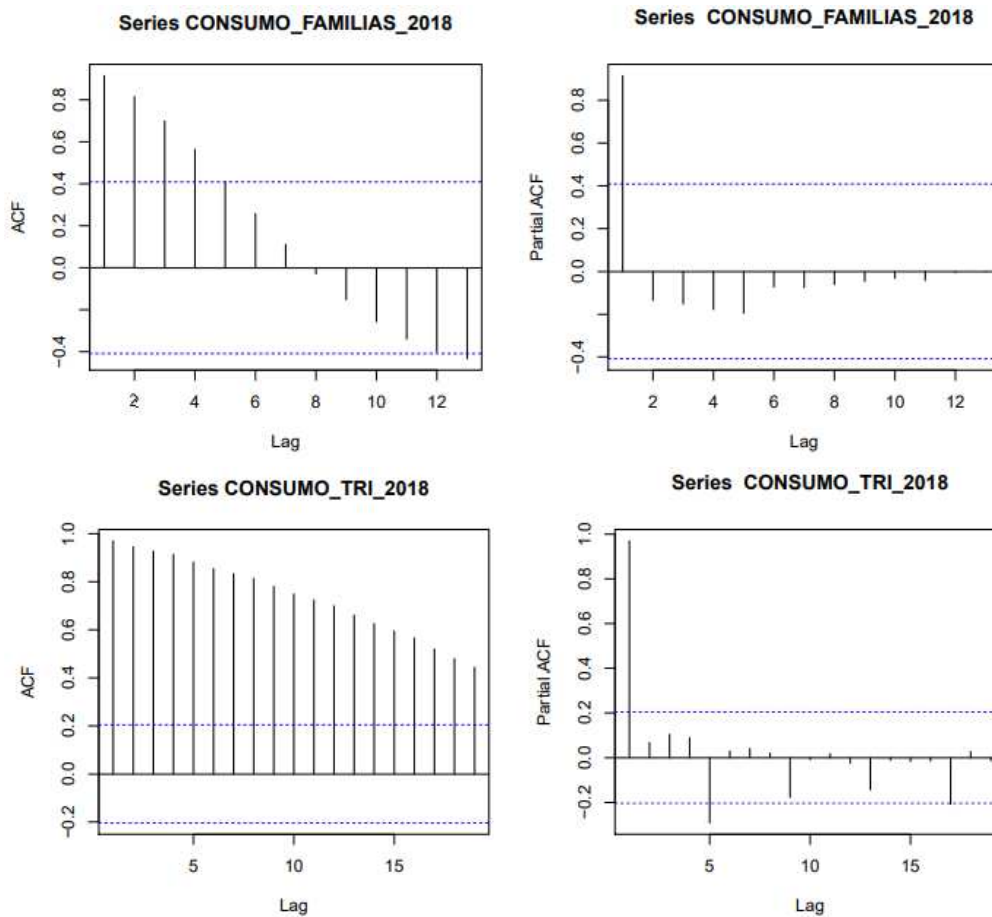
Figura 10: FAC e FACP do PIB anual de 1996 a 2018 / FAC e FACP do PIB trimestral de 1996.1 a 2018.4





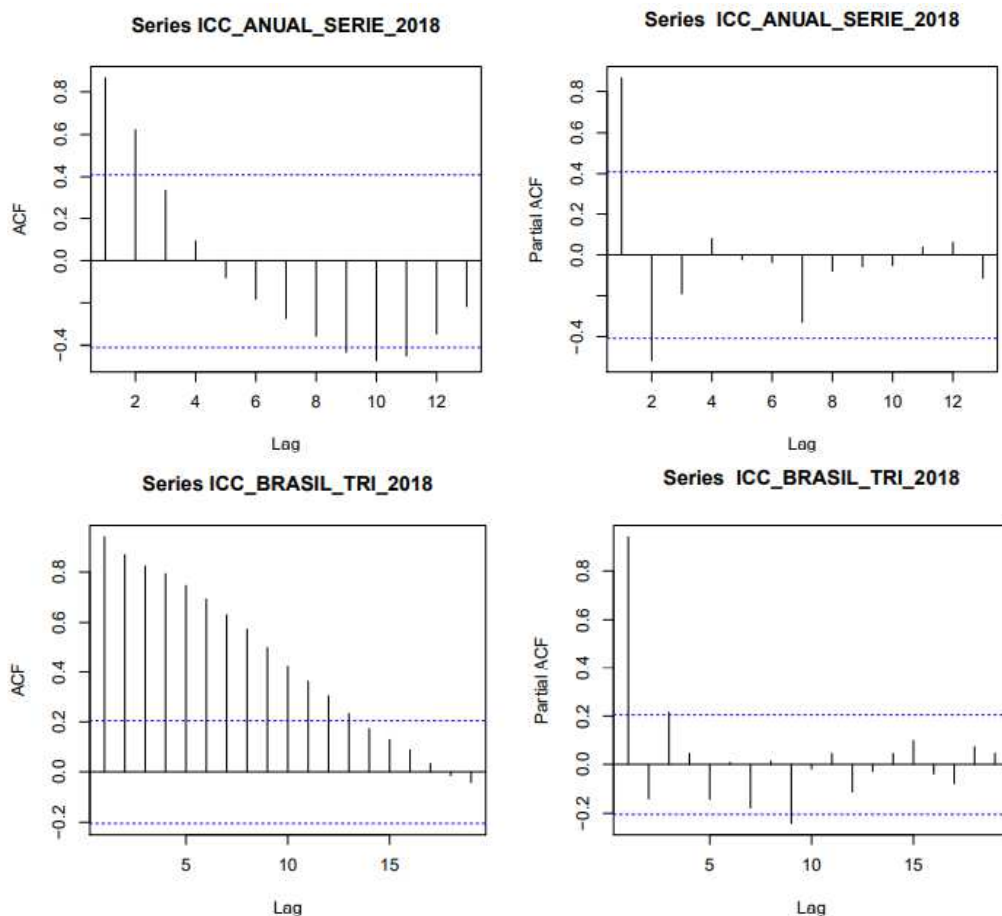
Fonte: O Autor

Figura 11: FAC e FACP do Consumo Privado anual de 1996 a 2018 / FAC e FACP do ICC
Consumo Privado trimestral de 1996.1 a 2018.4



Fonte: O Autor

Figura 12: FAC e FACP do ICC anual de 1996 a 2018 / FAC e FACP do ICC trimestral de 1996.1 a 2018.4



Fonte: O Autor

Utilizamos o teste de Philip-Perron para a presença ou não de raiz unitária, a hipótese nula do teste indica a presença de raiz unitária, presença de raiz unitária compromete a validação dos modelos econométricos, que são baseados em modelos que seguem uma média constante ao longo do tempo. Nas Tabelas 3 e 4 obteve-se os resultados do teste de Philip-Perron até cada série ter se tornado estacionária ao nível de confiança de 95%.

Em seguida, restava encontrar as diferenças sazonais das séries trimestrais, onde todas as séries ficaram sem sazonalidade após a primeira diferença sazonal.

Tabela 3: Teste de Philip-Perron após D diferenças das Séries Anuais e Trimestrais respectivamente, entre 1996-2017

P-Valor					
	Séries Anuais			Séries Trimestrais	
Série	D=0	D=1	D=2	D=0	D=1
PIB	0.7788	0.4962	0.01*	0.6971	0.01*
Consumo Privado	0.77077	0.6944	0.03985*	0.8511	0.01*
ICC	0.8772	0.3031	0.01*	0.8564	0.01*
p-valor significativo* < 0.05					

Fonte: O Autor

Tabela 4: Teste de Philip-Perron após D diferenças das Séries Anuais e Trimestrais respectivamente, entre 1996-2018

P-Valor					
	Séries Anuais			Séries Trimestrais	
Série	D=0	D=1	D=2	D=0	D=1
PIB	0.8371	0.4599	0.01*	0.7992	0.01*
Consumo Privado	0.7853	0.6574	0.0294*	0.8528	0.01*
ICC	0.8314	0.2785	0.01	0.8066	0.01*
p-valor significativo* < 0.05					

Fonte: O Autor

3.4 ESTIMAÇÃO DOS MODELOS DE SÉRIES TEMPORAIS

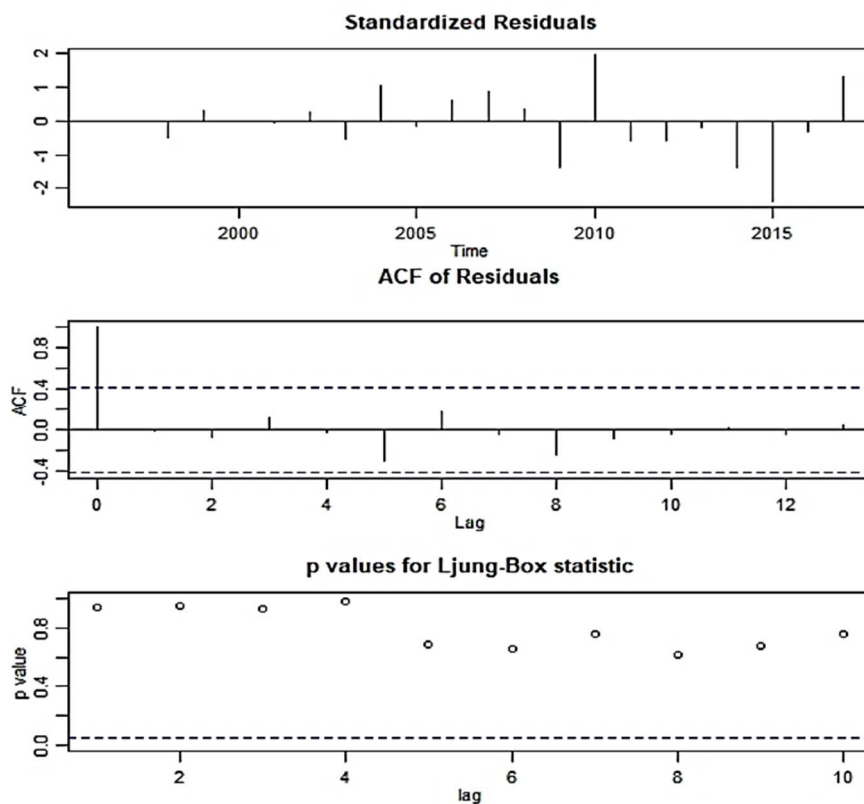
A função `auto.arima()`, é uma função utilizada no *software* R, que se utiliza de uma variação do algoritmo de Hyndman-Khandakar (Hyndman & Khandakar, 2008). Este algoritmo combina testes de raiz unitária, minimização do AICc e MLE afim de encontrar um modelo ARIMA.

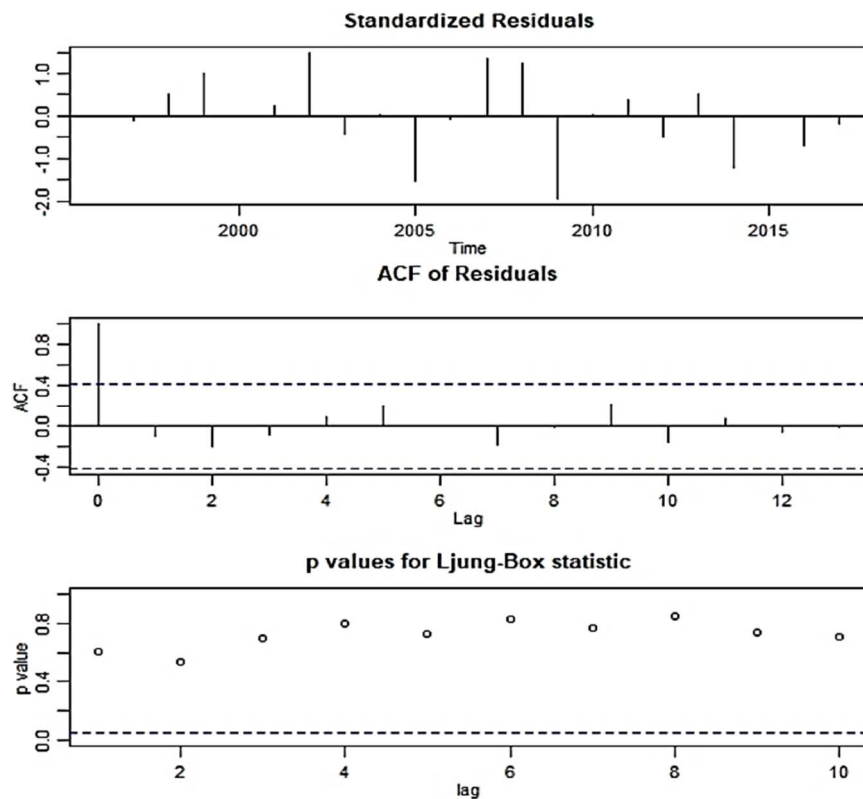
Na função `auto.arima`, podemos escolher quais os parâmetros dentro da própria função queremos usar, tal como os testes de estacionariedade e sazonalidade, assim como a inclusão de co-variáveis ao modelo, quando isso acontece, ou seja, quando há a inclusão de co-variáveis ao modelo dizemos que o modelo se trata de um ARIMAX, em que uma ou mais variáveis são adicionadas ao modelo afim de explicar juntamente com a parte da própria modelagem ARIMA a variação ao longo do tempo de Y_t .

Foi estimado primeiramente as séries anuais do PIB, por um modelo ARIMA(p,d,q), posteriormente o modelo ARIMAX(p,d,q) com as co-variáveis consumo privado anual e ICC anual. No caso trimestral, foram modeladas as séries do PIB trimestral, primeiro o modelo SARIMA(p,d,q)x(P,D,Q)[4], depois o modelo SARIMAX(p,d,q)x(P,D,Q)[4] com as co-variáveis trimestrais do Consumo privado e do ICC.

Todos os 8 modelos estimados apresentaram bom diagnóstico residual, tanto pelo ACF dos resíduos onde verificou-se a aderência dos mesmos ao processo ruído branco, quanto pelo teste de Box-Ljung onde foi verificado a auto-correlação entre os erros, habilitando-o a passar para a fase da previsão como mostram as Figuras 13, 14, 15 e 16. Um dos pressupostos das séries temporais estacionárias é que os resíduos sejam aleatórios, independentes e identicamente distribuídos e que se originem de um ruído branco.

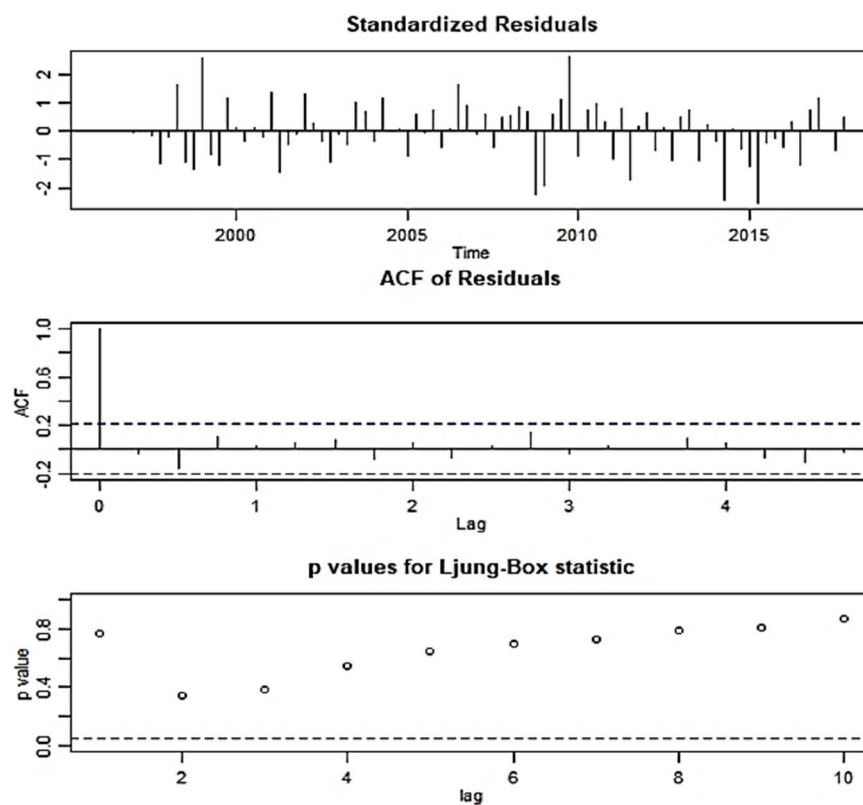
Figura 13: Diagnóstico Residual ARIMA(0, 2, 1) Série 1996-2017 / Diagnóstico Residual ARIMAX(2, 1, 0) Série 1996-2017

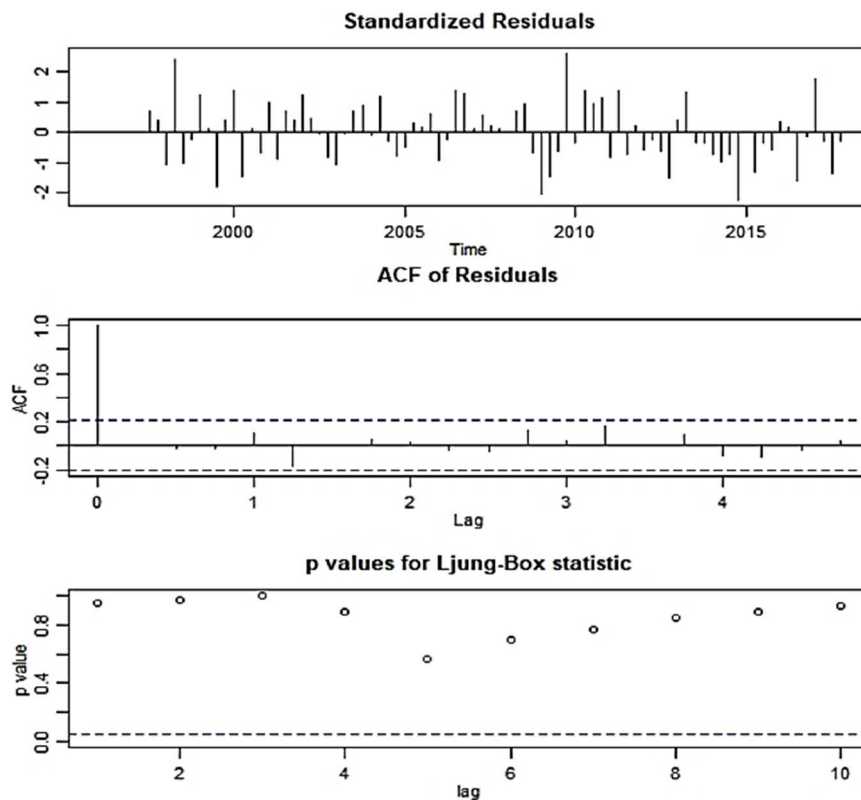




Fonte: O Autor

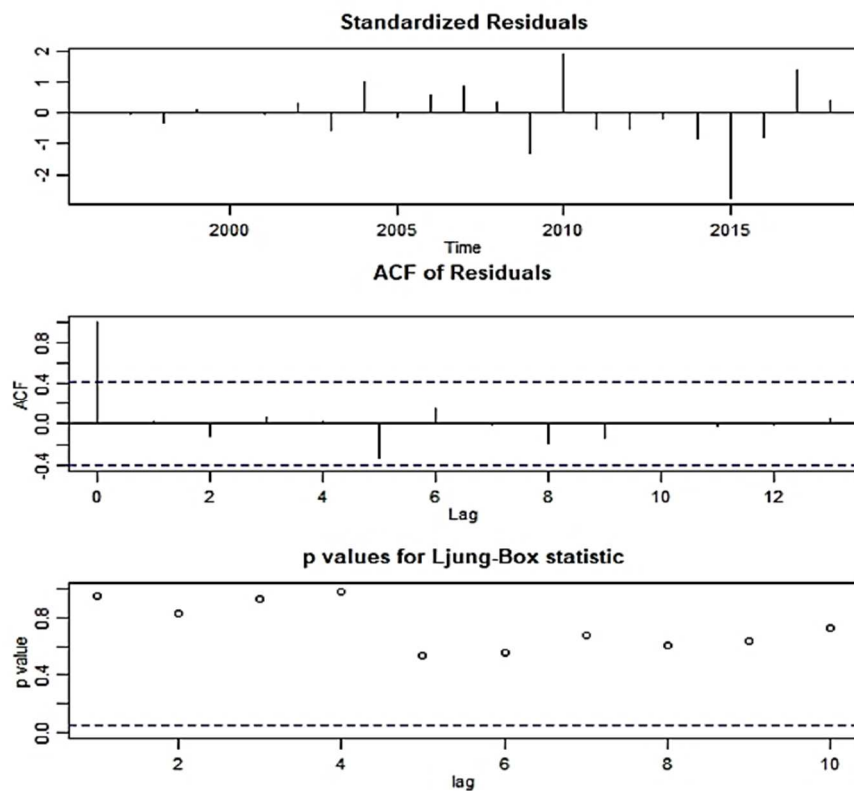
Figura 14: Diagnóstico Residual SARIMA(0, 1, 1)x(0, 1, 1) Série 1996.1-2017.4 / Diagnóstico Residual SARIMAX(0, 1, 2)x(0, 1, 1) Série 1996.1-2017.4

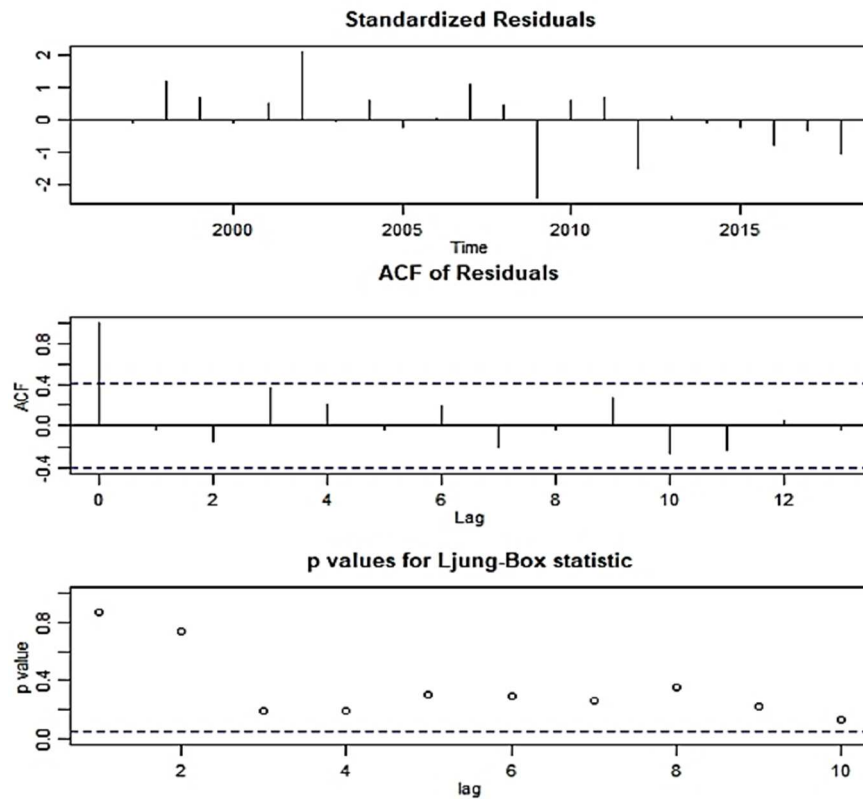




Fonte: O Autor

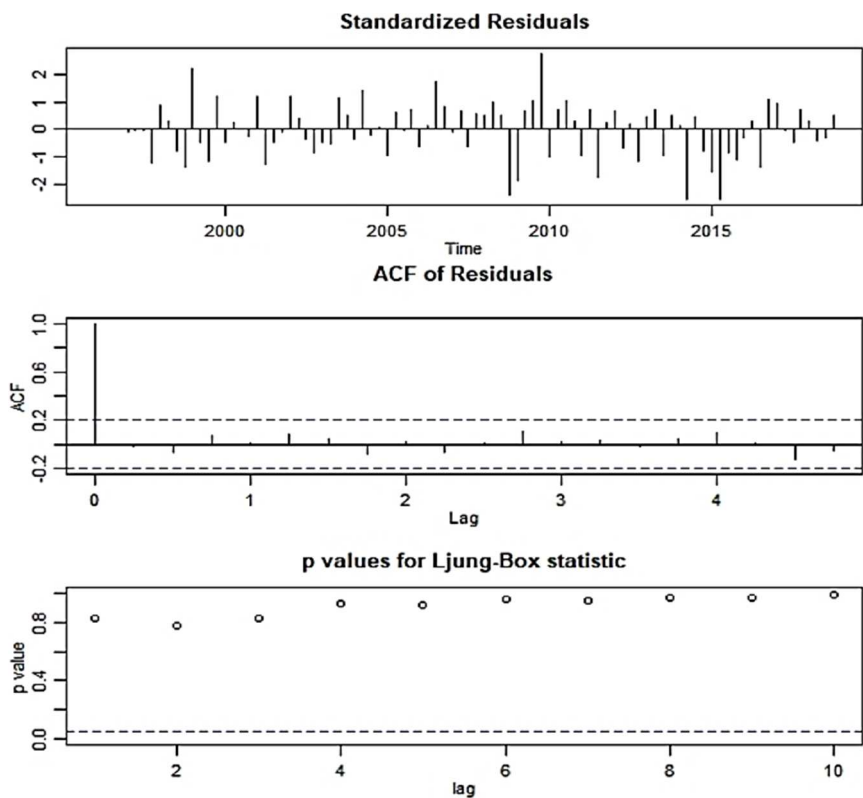
Figura 15: Diagnóstico Residual ARIMA(0, 2, 1) Série 1996-2018 / Diagnóstico Residual ARIMAX(0, 1, 1) Série 1996-2018

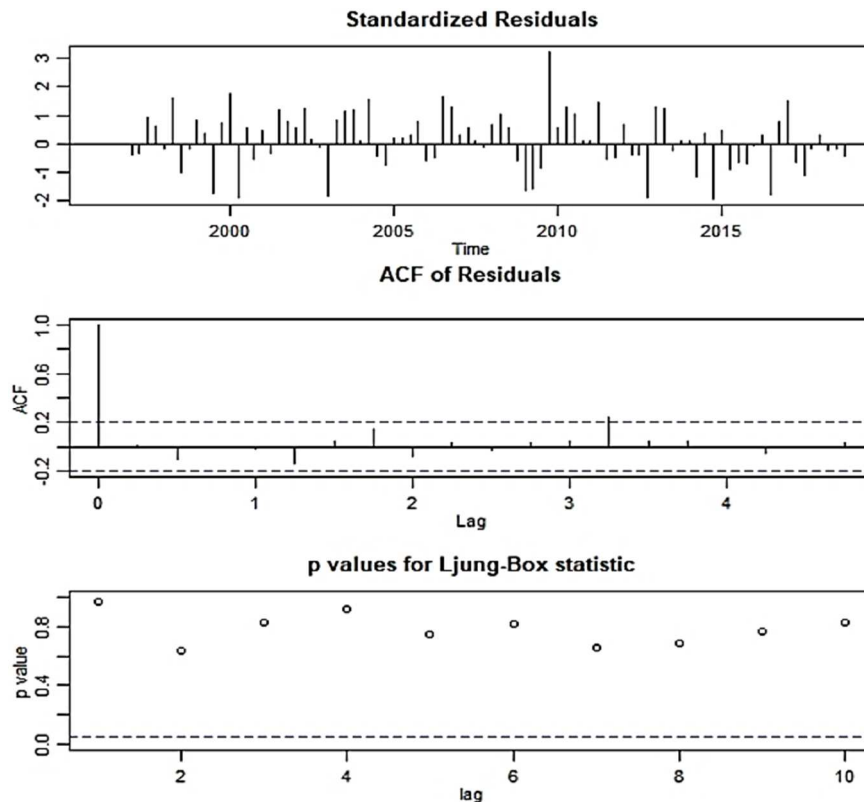




Fonte: O Autor

Figura 16: Diagnóstico Residual SARIMA(0, 1, 1)x(0, 1, 1)[4] Série 1996.1-2018.4 / Diagnóstico Residual SARIMAX(1, 0, 0)[4] Série 1996.1-2018.4

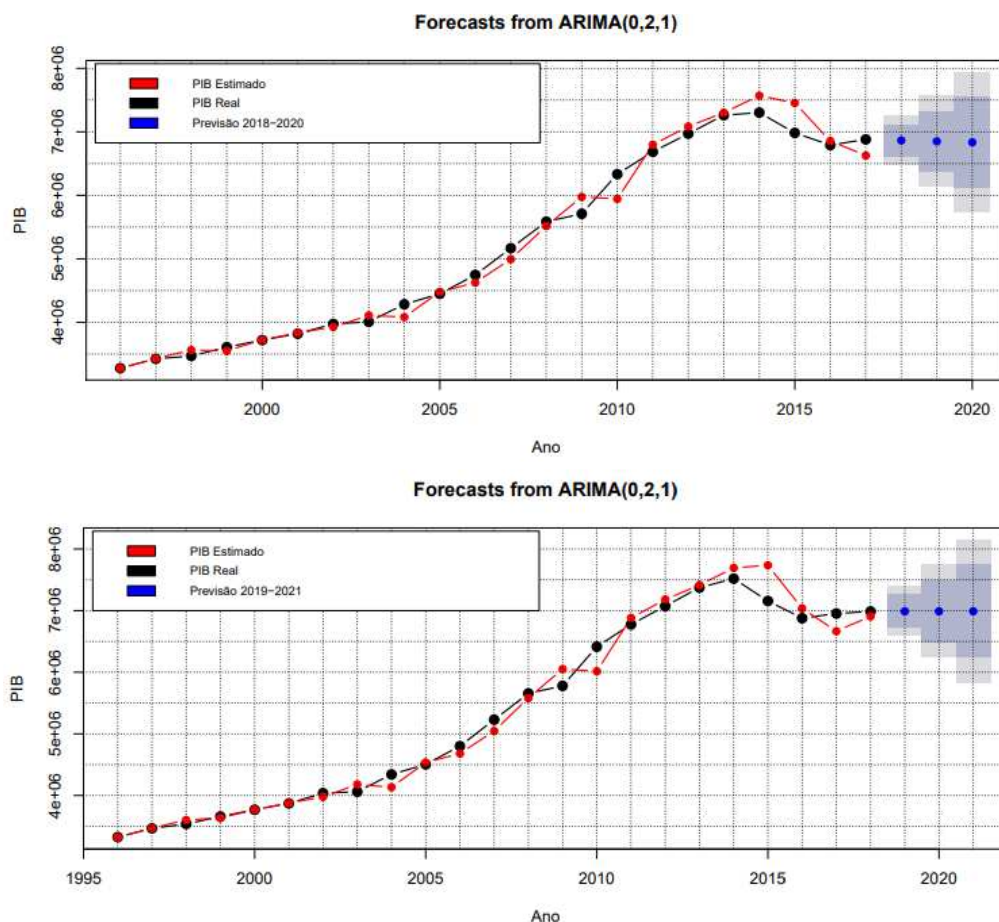




Fonte: O Autor

As previsões dos modelos anuais sem as co-variáveis com dados até 2017 e no período até 2018 mostraram-se pouco eficientes com intervalos de confiança cada vez maiores à medida que o tempo aumenta e não conseguindo captar uma tendência de recuperação econômica, assim como os próprios modelos que não se adequaram bem aos dados prevendo valores distorcidos sobretudo no período de 2014-2017 como mostra a Figura 17, onde os mesmos mostraram em 2015 um crescimento, não conseguindo captar a crise de 2014-2016, os mesmos modelos não previram um crescimento para 2017 invés disso segundo estes modelos teríamos tido outro ano de recessão, quando passamos para a previsão do modelo ARIMA(0,2,1) com dados até 2017 três anos à frente, claramente se notou que segundo este modelo teríamos mais três anos de recessão econômica, ou seja o modelo ARIMA(0,2,1) anual do PIB sem as co-variáveis não conseguiu captar o fim da recessão no ano de 2017 e nem prever crescimento entre 2018-2020, talvez pelo tamanho da amostra que não forneceu informações passadas suficientes ao modelo ou por não ter auxílio de outras variáveis. O modelo ARIMA(0,2,1) com dados atualizados de 2018 também prevê mais recessão até o ano de 2021, considerando uma previsão três anos à frente.

Figura 17: PIB Real versus o estimado pelo modelo ARIMA (0, 2, 1) Série 1996 – 2017 / PIB Real versus o estímulo pelo modelo ARIMAX (0, 2, 1) Série 1996 – 2018

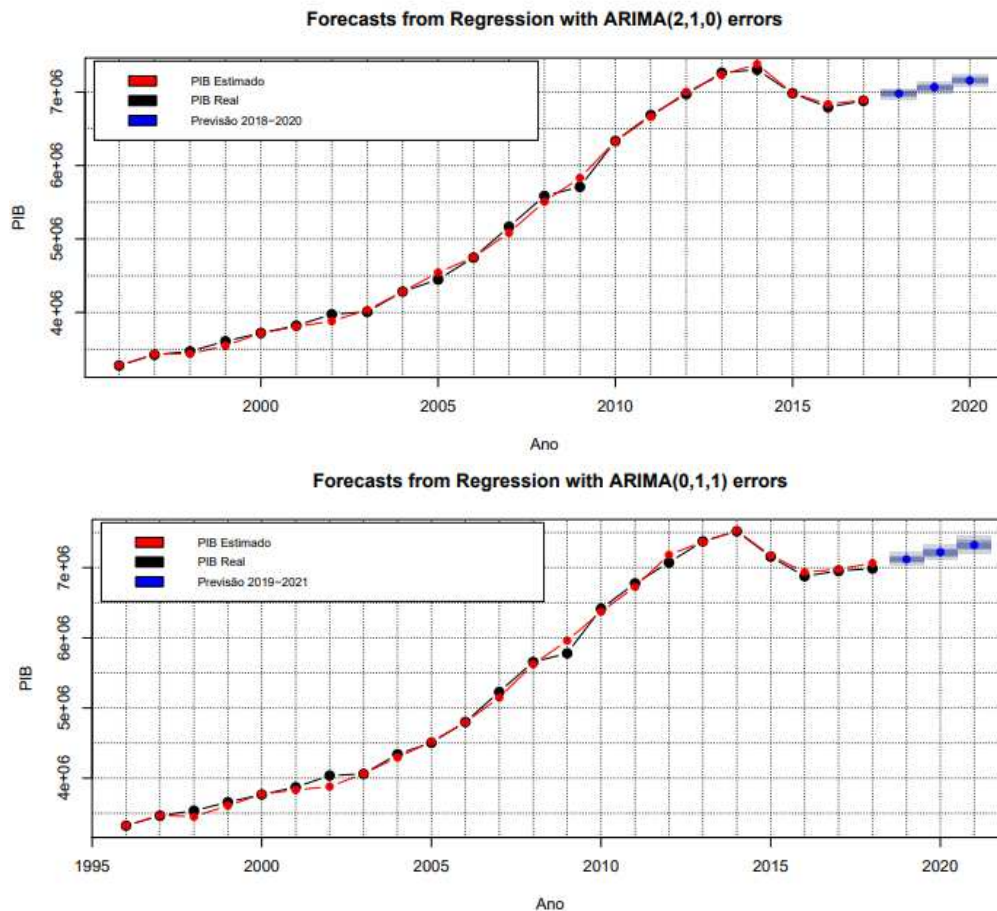


Fonte: O Autor

No modelo ARIMAX(2,1,0) com dados anuais até 2017 com as co-variáveis, notou-se um melhor ajuste do processo estimado com os valores reais do PIB anual (Figura 17), tendo sido os anos 2002 e 2009 os piores valores ajustados pelo modelo em comparação aos dados reais dos respectivos anos, o modelo se ajustou muito bem aos dados sobretudo a partir de 2013, conseguindo prever a recessão de 2014-2016 e a retomada econômica de 2017, sua previsão para 2018-2020 surpreendeu por mostrar uma recuperação lenta da economia, indicando que passaremos mais tempo para recuperar as perdas da depressão iniciada em 2014.

O modelo ARIMAX(0,1,1) com dados anuais no período entre 1996 e 2018, também prevê uma recuperação gradual da economia contudo, o ajuste parece ter ficado um pouco aquém se comparado com o ajuste do ARIMAX(2,1,0), como mostra a (Figura 18). Pode ser que isso tenha ocorrido pela atualização monetária e inflacionária calculadas no IPEADATA, contudo temos que verificar o EAM e o EQM para sermos mais conclusivos quanto a qualidade do ajuste.

Figura 18: PIB Real versus o estimado pelo modelo ARIMAX (2, 1, 0) Série 1996 – 2017 / PIB Real versus o estímulo pelo modelo ARIMAX (0, 1, 1) Série 1996 – 2018

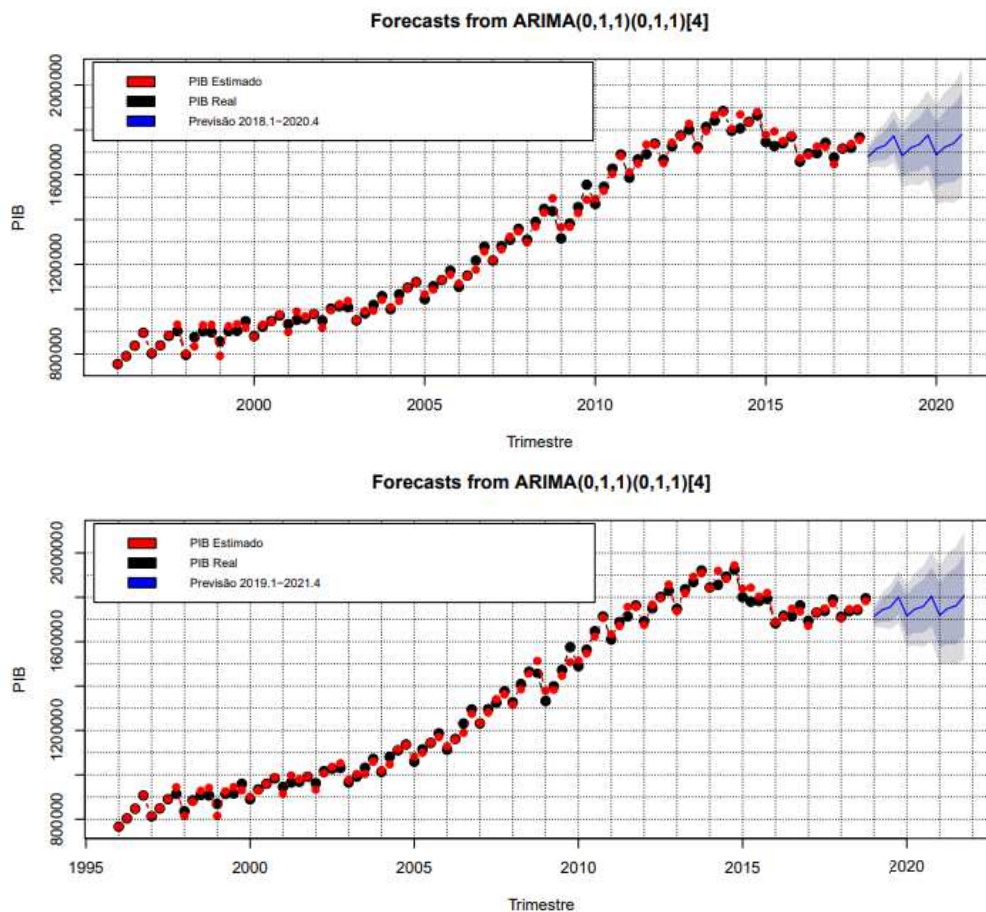


Fonte: O Autor

O modelo SARIMA(0,1,1)x(0,1,1)[4] da série trimestral do PIB no período de 1996.1 a 2017.4, conseguiu se adequar bem aos dados do PIB trimestral (Figura 18), conseguindo captar a crise econômica e a volta do crescimento em 2017, contudo o mesmo não conseguiu acompanhar com mais precisão as oscilações trimestrais do ano de 2015. A previsão para 2018.1-2020.4 segundo este modelo mostrou uma semi-estagnação do PIB, contudo seus intervalos de confiança são altos, supondo que o PIB se comporte sempre no limite superior do intervalo de 95% de confiança, teríamos então em 2020 recuperado o patamar de 2014, porém sabemos que nas condições atuais isso provavelmente é difícil de ocorrer.

Já o SARIMA(0,1,1)X(0,1,1)[4] da série trimestral do PIB entre 1996.1-2018.4, também prevê uma estagnação do produto brasileiro entre 2019-2021 (Figura 19), todavia seus intervalos de confiança são grandes, considerando que o PIB cresça no limite superior desse intervalo, teríamos em 2020 recuperado totalmente o Produto perdido com a recessão.

Figura 19: PIB Real versus o estimado pelo modelo SARIMA(0, 1, 1)x(0, 1, 1)[4] Série 1996.1 – 2017.4 / PIB Real versus o estímulo pelo modelo SARIMA(0, 1, 1)x(0, 1, 1)[4] Série 1996.1 – 2018.4



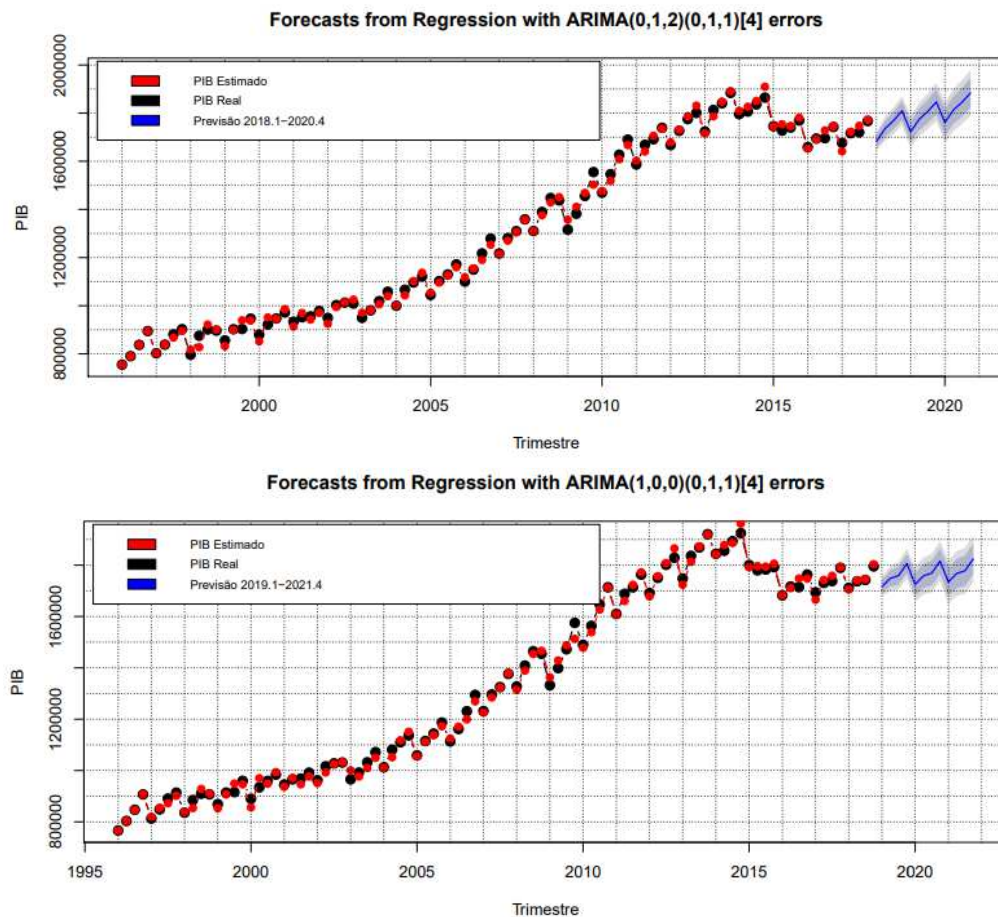
Fonte: O Autor

No modelo SARIMAX(0,1,2)x(0,1,1)[4] com dados trimestrais entre 1996.1-2017.4 com co-variáveis, os valores previstos também ficaram bem próximos do valor real como mostrou a Figura 19, um detalhe importante é que o SARIMAX(0,1,2)x(0,1,1)[4] conseguiu captar melhor os dados trimestrais de 2015 e de 2017, quanto a previsão, entre 2018.1-2020.4 o modelo previu uma recuperação gradual da economia trimestre a trimestre, os intervalos de confiança foram menores que o do SARIMA(0,1,1)x(0,1,1)[4], como no modelo anterior supondo que crescamos sempre no limite superior do intervalo a 95% confiança recuperaríamos em 2020 o patamar de 2014.

Já o SARIMAX(1,0,0)x(0,1,1)[4] ajustado sob os dados trimestrais do PIB entre 1996.1-2018.4 mostra uma recuperação ainda mais lenta se comparada a do SARIMAX(0,1,2)x(0,1,1)[4] com os dados terminado em 2017.4 (Figura 20), após a greve dos caminhoneiros o crescimento arrefeceu,

assim como o consumo das famílias e a confiança do consumidor, e segundo o mesmo modelo demoraremos mais tempo para voltar ao patamar de 2014, mesmo crescendo no limite superior do intervalo de confiança de 95%.

Figura 20: PIB Real versus o estimado pelo modelo SARIMAX(0, 1, 2)x(0, 1, 1)[4] Série 1996.1 – 2017.4 / PIB Real versus o estímulo pelo modelo SARIMAX(1, 0, 0)x(0, 1, 1)[4] Série 1996.1 – 2018.4



Fonte: O Autor

3.5 MEDIDAS DE ACURÁCIA

As Tabelas 5 e 6 mostram os respectivos valores para os dados anuais e trimestrais respectivamente comparando os modelos econométricos:

Tabela 5: Medidas de Acurácia, Modelos da Série 1996-2017

Modelo	Medidas	
	$\sqrt{EQM}$	EAM
ARIMA(0, 2, 1)	181852.2	131564.7
ARIMAX(2, 1, 0)	53851.41	37993.52
SARIMA(0, 1, 1)x(0, 1, 1) [4]	24610.76	18701.06
SARIMAX(0, 1, 2)x(0, 1, 1) [4]	18892.34	14697.12

Fonte: O Autor

Analisando a Tabela 5, vemos que os modelos ARIMAX possuem menor desvio do que ao modelo ARIMA convencional. Além disso. Inserindo sazonalidade nestes modelos também tornam melhores as medidas de ajuste. Resultado semelhante se observa na Tabela 6 para os dados trimestrais. Assim, podemos supor que modelos sazonais que levam em consideração covariáveis serão aqueles que poderão possuir previsões mais próximas a realidade para os dados futuros.

Tabela 6: Medidas de Acurácia, Modelos da Série 1996-2018

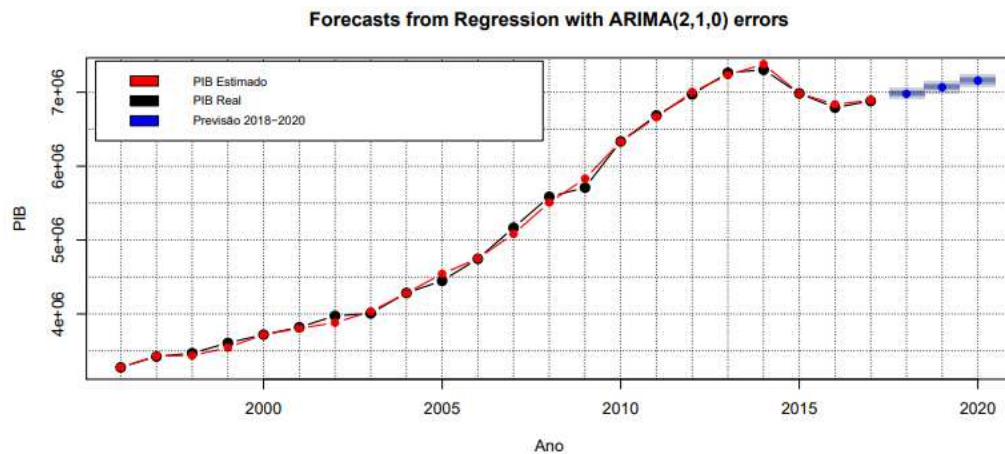
Modelo	Medidas	
	$\sqrt{EQM}$	EAM
ARIMA(0, 2, 1)	193101.3	134903.2
ARIMAX(1, 1, 0)	69104.63	47803.63
SARIMA(0, 1, 1)x(0, 1, 1) [4]	23600.34	18166.07
SARIMAX(1, 0, 0)x(0, 1, 1) [4]	18252.26	14139.83

Fonte: O Autor

3.6 PREVISÃO

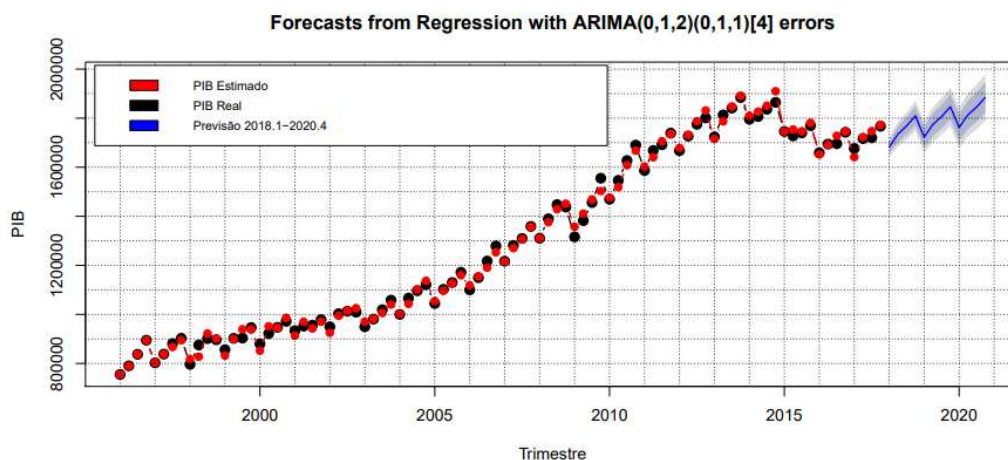
Passadas todas as etapas de diagnósticos foram feitas as previsões dos modelos entre 2018-2020, como mostram as Figuras 21, 22, 23 e 24.

Figura 21: Previsão do PIB brasileiro entre 2018-2020 segundo Modelo ARIMAX(2, 1, 0)



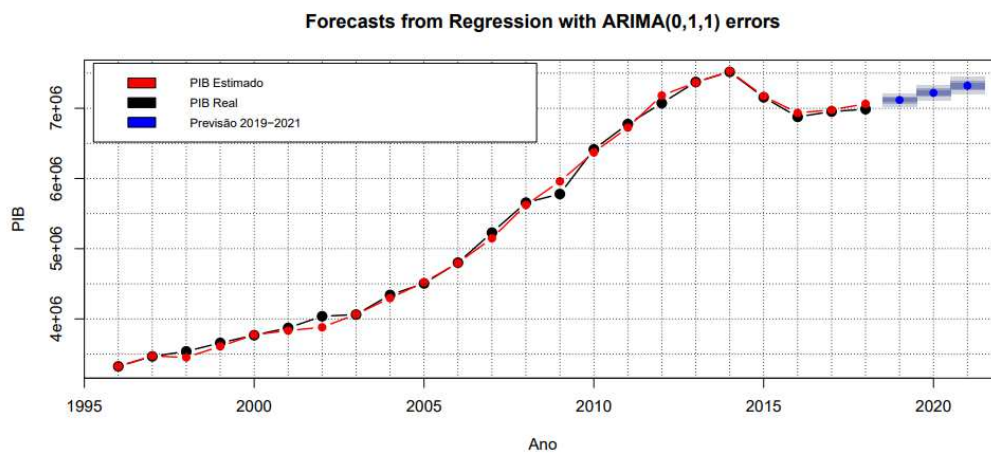
Fonte: O Autor

Figura 22: Previsão do PIB brasileiro entre 2018.1-2020.4 segundo Modelo SARIMAX(0, 1, 2)x(0, 1, 1)[4]



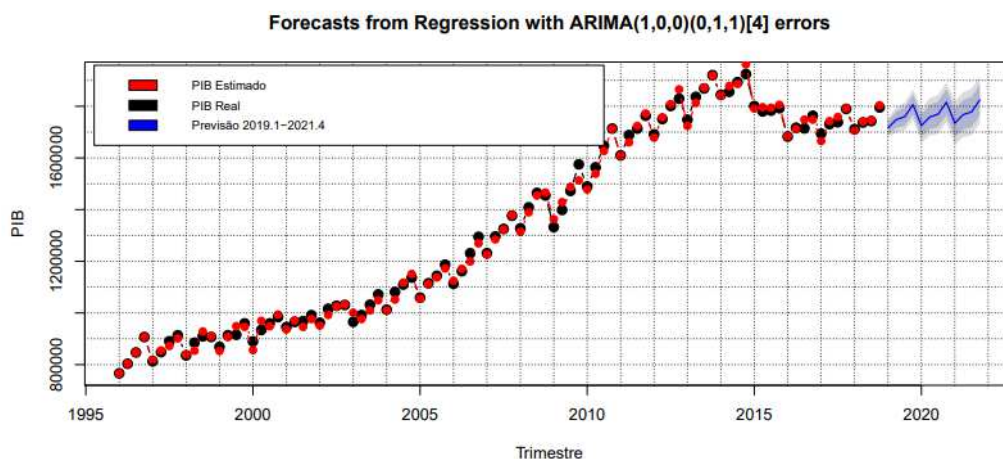
Fonte: O Autor

Figura 23: Previsão do PIB brasileiro entre 2019-2021 segundo Modelo ARIMAX(0, 1, 1)



Fonte: O Autor

Figura 24: Previsão do PIB brasileiro entre 2019.1-2021.4 segundo Modelo SARIMAX(1, 0, 0)x(0, 1, 1)[4]



Fonte: O Autor

As Figuras 21 e 22 mostraram uma previsão de recuperação gradual do PIB brasileiro entre 2018-2020, vindo de encontro ao que os analistas vêm falando a respeito da economia brasileira, que ela irá demorar para se recuperar da recessão recente. No modelo da Figura 21, mesmo considerando o limite superior do intervalo a 95% de confiança em 2020 ainda não teríamos recuperado o patamar de 2014, contudo na Figura 22 se consideramos o limite superior do intervalo a 95% de confiança teremos atingido no último trimestre de 2020, um valor muito próximo ao do último trimestre de 2014, a Tabela 7 mostra a comparação entre o valor real e o previsto pelo modelo SARIMAX(0,1,2)x(0,1,1) para os trimestres de 2018, podemos notar que o modelo erra pouco no primeiro trimestre e no terceiro trimestre de 2018, a greve dos caminhoneiros teve um papel importante na redução do ritmo de crescimento econômico e isso contribuiu para que o modelo errasse mais na

previsão de 2018.2. No ano de 2018 o PIB teve um crescimento real de 1.1% com relação ao ano anterior, e o modelo ARIMAX(2,1,0) (Figura 21), conseguiu prever um crescimento mais próximo da taxa de crescimento do produto se comparada à previsão de crescimento que as instituições oficiais e extra-oficiais previam, considerando que os dados coletados para essa pesquisa foram consolidados em março de 2018, antes da greve dos caminhoneiros, nesse mesmo sentido o SARIMAX(0,1,2)X(0,1,1), como mostra a Tabela 9 também se aproximou da previsão real do PIB em 2018.

Tabela 7: PIB Trimestral versus PIB Estimado pelo SARIMAX(0, 2, 1)x(0, 1, 1)[4] para 2018

Trimestre	PIB		
	Real	Previsto	Diferença (%) [ii/iii]
2018.1	1688985.48	1683558	0.32
2018.2	1720092.39	1734628	0.84
2018.3	1764835.96	1768510	0.21
2018.4	*	1810822	-

Fonte: O Autor

Tabela 8: Previsões de Crescimento Anual do PIB do Brasil para o ano de 2018 divulgadas no fim de 2018, comparadas com as estimativas dos modelos ARIMAX e SARIMAX

Fontes	PIB
	Previsão de Crescimento (%)
FMI (Outubro/2018)	1.4
Banco Mundial (Outubro/2018)	1.2
Boletim Focus (Novembro/2018)	1.39
Ministério da Fazenda (Novembro/2018)	1.4
FGV (Novembro/2018)	1.5
OCDE (Outubro/2018)	1.2
Modelo ARIMAX (2, 1, 0)	1.42
Modelo SARIMAX (0, 2, 1)x(0, 1, 1)[4]	1.7

Fonte: O Autor

Tabela 9: Previsões de Crescimento Anual do PIB do Brasil para o ano de 2018 divulgadas antes da Greve dos Caminhoneiros e as previsões dos Modelos ARIMAX e SARIMAX

Fontes	PIB
	Previsão de Crescimento (%)
FMI (Abril/2018)	2.3
IPEA (Março/2018)	3
Boletim Focus (Abril/2018)	2.75
Modelo ARIMAX (2, 1, 0)	1.42
Modelo SARIMAX (0, 2, 1)x(0, 1, 1)[4]	1.7

Fonte: O Autor

Tabela 10: Previsão de Crescimento Anual do PIB do Brasil para o ano de 2019

	PIB
Fontes	Previsão de Crescimento (%)
FMI (Outubro/2018)	2.4
Boletim Focus (Novembro/2018)	2.5
Modelo ARIMAX (2, 1, 0)	1.27
Modelo SARIMAX (0, 2, 1)x(0, 1, 1)[4]	2.22

Fonte: O Autor

Ao compararmos as previsões de crescimento para o PIB segundo os modelos estimados por esse estudo com a previsão de outras fontes notamos que os modelos preveem bem o crescimento uma vez que suas estimativas se baseiam entre 1996-2017 e mesmo faltando as informações parciais já fornecidas nesse ano sobre o produto conseguiram resultados satisfatórios como demonstrou a Tabela 8, o mais correto foi comparar as estimativas dos modelos com as estimações das instituições econômicas no começo de 2018, como mostrou a Tabela 9. As estimativas realizadas no começo deste ano por estas instituições dista muito das previsões dos dias atuais, os modelos estimados aqui nesse estudo se saem melhor e são os que mais chegaram perto do que se prevê de crescimento para esse ano e com o detalhe de estar usando os mesmos dados que eram disponíveis até então.

A estimativa para 2019 do modelo SARIMAX(0,1,2)x(0,1,1)[4] se mostrou próxima as estimativas das outras fontes, já o modelo ARIMAX(2,1,0) não conseguiu prever uma melhora na expectativa de crescimento econômico.

4. CONCLUSÃO

Os modelos ARIMAX e SARIMAX conseguiram explicar melhor a variação do PIB tanto anual quanto trimestralmente, reduzindo seu erro médio e oferecendo estimativas precisas. O ARIMAX mostrou que nem sempre precisamos de uma amostra grande para podermos explicar a variação da variável de interesse, escolhendo as co-variáveis certas podemos ter resultados surpreendentes. O SARIMAX também se mostrou eficiente pois a presença co-variáveis foi de primordial importância na redução do erro médio e na previsão t passos a frente.

As previsões para 2019 vão desde uma recuperação gradual como vemos no caso do SARIMAX(0,2,1)x(0,1,1)[4] e ARIMAX(0,1,1) e que vai de encontro com a maioria da expectativas atuais, quanto uma recuperação lenta como nos casos do ARIMAX(2,1,0) e do SARIMAX(1,0,0)x(0,1,1)[4]. A greve dos caminhoneiros e as incertezas políticas e econômicas levaram a uma piora do crescimento econômico e da confiança do consumidor, o

SARIMAX(1,0,0)x(0,1,1) prevê que essa incerteza e a continuidade desse modelo de política econômica pode fazer o país ter um crescimento pequeno por mais um ano, indicando a severidade da depressão econômica e a dificuldade do país recuperar taxas sólidas de crescimento econômico.

5. REFERÊNCIAS BIBLIOGRÁFICAS

Aranha, D. M. (2017). *A relação dos indicadores de confiança com o crescimento econômico*. 2017. 45f. Dissertação de Mestrado em Economia - FGV São Paulo, São Paulo, 2017.

Breusch, T. S. & Pagan, A. R. (1979). *A Simple Test for Heteroskedasticity and Random Coefficient Variation*. *Econometrica*. 1287–1294.

Charnet, R. *et al.* (2015). *Análises de Modelos de Regressão Linear: Com aplicações*. 2ª Edição, 5–296. ~ Editora Unicamp. Campinas: São Paulo, 2015.

Chan, K. S. & Ripley, B. (2018). *TSA: Time Series Analysis*. R package version 1.2. <https://CRAN.Rproject.org/package=TSA>

Dados Abertos Banco Central (2018). *Índice de Confiança do Consumidor Paulistano*. Disponível em: <https://dadosabertos.bcb.gov.br>, Acessado em: 20 de Novembro de 2018.

Durbin, J. & Watson, G. S. (1950). *Testing for Serial Correlation in Least Squares Regression*, I. *Biometrika*. 409–428.

Emiliano, P. C. (2009). *Fundamentos e Aplicações dos Critérios de Informação: Akaike e Bayesiano*. 2009. 105f. Dissertação de Mestrado em Estatística e Experimentação Agropecuária - UFPA, Belém, 2009.

Fecomércio – SP (2018). *Índice de Confiança do Consumidor*. São Paulo: Fecomércio - SP, 2018. Disponível em: <http://www.fecomercio.com.br/pesquisas/indice/icc>, Acessado em: 29 de Novembro de 2018.

Fischer, S. (1982). *Series univariantes de tempo — metodologia de Box & Jenkins*. Porto Alegre, FEE.

Forecasting: Principles and Practices (2018). *Modeling ARIMA in R - How does auto.arima() work?*. Disponível em: <https://otexts.org/fpp2/arima-r.html>, Acessado em: 25 de Novembro de 2018.

Gerolimetto, M. & Procidano, I. (2008). *Regime-sensitive cointegration: the relationship between gold and silver*. *Far East Journal of Theoretical Statistics*, vol. 30. 41-47. Disponível em: <https://www.unive.it/data/persona>, Acessado em: 29 de Novembro de 2018.

Gomes, O. (2012). *Macroeconomia: Noções Básicas*. Disponível em: <https://repositorio.ipl.pt/bitstream/10400.21/1186/1/MacroIntroCap.pdf>, acessado em: 29 de Novembro de 2018.

Gurajati, D. N. & Porter, D. C. (2011). *Econometria Básica*. 5ª edição, Nova Iorque: Editora AMGH, 2011.

- Horthron, T. & Zeileis, A. & Farebrother, R. W. & Cummins, C. & Millo, G. & Mitchell, D. (2019). *lmtest: Testing Linear Regression Models*. R package version 0.9-37. Disponível em: <https://CRAN.R-project.org/package=lmtest>.
- Hyndman, R. J. & Khandakar, Y. (2008). *Automatic Time Series Forecasting: The Forecast Package for R*. *Journal of statistical software*. volume 27. 1–22.
- IPEADATA (2018). *Macroeconomico – PIB*. Disponível em: <https://www.ipeadata.gov.br>, Acessado em: 29 de Novembro de 2018.
- IPEADATA (2018). *Macroeconomico - Consumo Familiar Final*. Disponível em: <https://www.ipeadata.gov.br>, Acessado em: 29 de Novembro de 2018.
- Ljung, G. M. & Box, G. E. P. (1978). *On a measure of a Lack of Fit in Time Series Models*. *Biometrika*. 297–303.
- Morettin, P. A. & Toloi, C. M. C. (2006). *Análise de Séries temporais*. 2ª edição. São Paulo: Editora Blucher.
- Pfaff, B. & Zivot, E. & Stigler, M. (2016). *URCA: Unit Root and Cointegration Tests for Time Series Data*. R package version 1.3-0. Disponível em: <https://CRAN.R-project.org/package=urca>.
- Phillips, P. C. B. & Perron, P. (1988). *Testing for a Unit Root in Time Series Regression*. 335–346.
- Portal Action (2018). *Teste de Phillips*. São Paulo: Portal Action. Disponível em: <http://www.portalaction.com.br/series-temporais/142-teste-de-phillips-perron>, Acessado em: 10 de Novembro de 2018.
- Portal Action (2018). *Teste de Ljung-Box*. São Paulo: Portal Action. Disponível em: <http://www.portalaction.com.br/series-temporais/472-teste-de-ljung-box>, Acessado em: 10 de Novembro de 2018.
- Shapiro, S. S. & Wilk, M. B. (1965). *An analysis of variance test for normality (complete samples)*. *Biometrika*. 591–611.
- R Core Team – R (2018). *A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria. 591–611. Disponível em: <https://www.R-project.org/>.
- Silva, A. H. A. (2013). *Testes escore para transformação de dados em regressões lineares*. 108f. Dissertação de Mestrado em Estatística - UFPE, Recife, 2013. Disponível em: <https://repositorio.ufpe.br/bitstream>, Acessado em: 14 de Dezembro de 2018.
- Trapletti, A. & Hornik, K. & Lebaron, B. (2019). *tseries: Time Series Analysis and Computational Finance*. R package version 0.10-47. <https://CRAN.R-project.org/package=tseries>.

ISSN 2675-3243

volume 78

número 243

janeiro/junho 2020