



Universidade Federal da Paraíba

Centro de Tecnologia

PROGRAMA DE PÓS-GRADUAÇÃO EM ENGENHARIA CIVIL E

AMBIENTAL

– MESTRADO/DOUTORADO –

**PREDIÇÃO DA PRECIPITAÇÃO NA AMAZÔNIA LEGAL USANDO
MODELOS ADITIVOS GENERALIZADOS DE LOCAÇÃO, ESCALA
E FORMA (GAMLSS)**

Por

Raul Souza Muniz

*Dissertação de Mestrado apresentada à Universidade Federal da Paraíba
para obtenção do grau de Mestre*

João Pessoa – Paraíba

Julho de 2024



Universidade Federal da Paraíba

Centro de Tecnologia

**PROGRAMA DE PÓS-GRADUAÇÃO EM ENGENHARIA CIVIL E
AMBIENTAL
– MESTRADO/DOCTORADO –**

**PREDIÇÃO DA PRECIPITAÇÃO NA AMAZÔNIA LEGAL USANDO
MODELOS ADITIVOS GENERALIZADOS DE LOCAÇÃO, ESCALA
E FORMA (GAMLSS)**

Dissertação submetida ao Programa de Pós-Graduação em Engenharia Civil e Ambiental da Universidade Federal da Paraíba, como parte dos requisitos para a obtenção do título de Mestre.

Raul Souza Muniz

Orientador: Prof. Dr. Celso Augusto Guimarães Santos

Coorientadora: Profa. Dra. Fernanda De Bastiani

João Pessoa – Paraíba

Julho de 2024

Catálogo na publicação
Seção de Catalogação e Classificação

M966p Muniz, Raul Souza.

Predição da precipitação na Amazônia legal usando
modelos aditivos generalizados de locação, escala e
forma (GAMLSS) / Raul Souza Muniz. - João Pessoa, 2024.
166 f. : il.

Orientação: Celso Augusto Guimarães Santos.

Coorientação: Fernanda de Bastiani.

Dissertação (Mestrado) - UFPB/CT.

1. Precipitação. 2. Teleconexões. 3. Amazônia legal.
4. GAMLSS. I. Santos, Celso Augusto Guimarães. II.
Bastiani, Fernanda de. III. Título.

UFPB/BC

CDU 502 (811.3) (043)

**PREDIÇÃO DA PRECIPITAÇÃO NA AMAZÔNIA LEGAL USANDO
MODELOS ADITIVOS GENERALIZADOS DE LOCAÇÃO, ESCALA E
FORMA (GAMLSS)**

RAUL SOUZA MUNIZ

Dissertação aprovada em 29 de julho de 2024.

Período Letivo: 2024.1

Documento assinado digitalmente



CELSO AUGUSTO GUIMARAES SANTOS
Data: 04/08/2024 23:30:03-0300
Verifique em <https://validar.iti.gov.br>

Profa. Dra. Celso Augusto Guimarães Santos – UFPB
Orientador

Documento assinado digitalmente



FERNANDA DE BASTIANI
Data: 05/08/2024 06:42:25-0300
Verifique em <https://validar.iti.gov.br>

Profa. Dra. Fernanda De Bastiani – UFPE
Coorientadora

Documento assinado digitalmente



RICHARDE MARQUES DA SILVA
Data: 05/08/2024 01:04:22-0300
Verifique em <https://validar.iti.gov.br>

Prof. Dr. Richarde Marques da Silva – UFPB
Examinador Interno

Gabriel de Oliveira

Prof. Dr. Gabriel de Oliveira – USA
Examinador Externo

João Pessoa/PB
2024

Dedico este trabalho aos meus familiares e a todas as pessoas que contribuíram para a minha jornada até aqui. A vocês, minha eterna gratidão.

Posso, tudo posso, naquele que me fortalece.
Nada e ninguém no mundo vai me fazer
desistir.

AGRADECIMENTOS

Agradeço primeiramente a Deus, por todas as bênçãos derramadas ao longo da minha jornada, por ter me dado esperança e força quando eu não tinha nenhuma, por ter me levantado quando eu estive caído na praia, e por ter me concedido todas os milagres até aqui vividos por mim.

Sou imensamente grato aos meus pais, Artur e Cláudia, pelo amor incondicional ao longo da minha vida, por sempre colocarem meu bem-estar em primeiro lugar, priorizarem minha educação e me apoiarem em todas as ocasiões. Graças à educação baseada no incentivo aos estudos que eles me proporcionaram, consegui persistir e progredir no meio acadêmico. Se não fosse por eles, possivelmente, não estaria alcançando estes avanços na minha jornada acadêmica. À minha irmã Aíla, minha eterna parceira, agradeço pelos conselhos, apoio e companheirismo, e por todas as nossas loucuras que tornam a vida mais leve e nos ajudam a enfrentar os desafios do mundo. Vocês são, sem dúvida, as melhores pessoas que eu poderia ter ao meu lado. Esta vitória não é apenas minha, é nossa.

Agradeço a minha namorada Laís. Ao longo desses anos, muita coisa mudou, mas sua preocupação com meu futuro sempre permaneceu constante. Muito obrigado por tudo, sou imensamente grato a Deus por ter te encontrado.

Vidinha sempre foi minha melhor amiga, nos momentos mais obscuros da minha vida foi capaz de tirar sorrisos do meu rosto e melhorar os meus dias. Hoje, Vidinha não está mais entre nós fisicamente, contudo, nos dias em que os scripts não rodavam e parecia que a dissertação não ia sair, a simples lembrança dela me ajudava a manter o foco.

Agradeço a minha avó Rosalia, que está festejando hoje lá do céu essa conquista, e a minha avó Socorro, cuja representatividade na minha vida é inestimável. Ambas sempre me incentivaram para que eu alcançasse o sucesso profissional.

Agradeço aos meus amigos do Banco do Brasil: Daniel Cabral, Derick, Gorito, Bárbara, Carol, Daniel Aires, Bruno Amorim, Pedro e Germano, assim como aos PadaUAns: Ramon, Paulo, Cadu, Vanine, Camila, Marcos, Quintella e Fernando. A mudança de cidade, estado, região, trabalho, clima e vida que experimentei ao longo desta jornada do mestrado não foi fácil, e as decisões que precisei tomar foram ainda mais difíceis. Entretanto, o acolhimento que recebi foi indescritível e incomparável; me senti abraçado e apoiado por todos. Embora alguns nem soubessem, vocês me ajudaram imensamente em muitos dias difíceis.

Gostaria de destacar a participação de Daniel Cabral, que me incentivou a ingressar no Banco do Brasil e embarcou comigo nessa mudança de vida. Ele abriu as portas para que eu conhecesse seus maravilhosos familiares: André, Débora, Cabralzinho e Clarissa. Nossa parceria, que começou na graduação e foi intensificada pelo período em que moramos juntos na Espanha, só tem se fortalecido. Espero que o destino continue a nos manter sempre próximos.

Gostaria de agradecer a Toscano que, com sua sanfona, sempre traz diversão garantida cada vez que nos encontramos, me ajudando a superar as dificuldades na construção deste trabalho.

Agradeço aos amigos Deividy, Mikael e Gustavo que compartilharam a sala de aula comigo durante o mestrado e sempre me ajudaram muito.

Agradeço ao meu professor e orientador Dr. Celso Augusto, ao qual tenho imensa admiração como profissional e como pessoa. Foi por meio do nosso segundo encontro na UFPB que resolvi migrar para área de recursos hídricos e abraçar a programação como ferramenta de trabalho. Seus incentivos, conselhos e correções me auxiliam a crescer cada dia mais, espero que possa continuar sempre usufruindo de sua experiência.

Agradeço à professora Fernanda, cuja contribuição foi fundamental desde o primeiro dia em que decidimos seguir com o tema deste trabalho. Ela sempre se mostrou muito disposta a ajudar e aceitou prontamente o desafio de ser minha coorientadora.

Agradeço a Leydson pela grande parceria e disponibilidade em compartilhar todo o seu conhecimento acerca do GAMLSS e pelo seu belíssimo trabalho. Agradeço igualmente ao professor Carlos Antonio, da UFCG, que tive o prazer de ter na minha banca de TCC e na qualificação, e ao professor Richarde Marques, que esteve presente na minha banca de TCC, na minha qualificação e também compõe a banca na minha defesa da dissertação. Com seus valiosos comentários e propostas de melhoria, tornaram este trabalho cada vez melhor. Agradeço também ao professor Gabriel de Oliveira por aceitar o convite de compor minha banca e nos auxiliar na construção deste trabalho. Sou grato aos professores Cristiano e Gerald por todo o conhecimento ensinado no âmbito do geoprocessamento e processamento em nuvens, e pelo privilégio de ter compartilhado uma sala de aula com profissionais de tamanha grandeza.

Agradeço a Felipe Alves, a Thiago Nascimento, a Abner Lins, a Deivyson e a Gabriel Rairan pelos conselhos, discussões, por compartilhar as dificuldades e por torcer por essa

conquista. Apesar de estar muito distante fisicamente deles, sempre me sinto muito próximo. Obrigado por estarem sempre comigo.

Sou grato a Universidade Federal da Paraíba por essa oportunidade, ao CNPq e a CAPES que me proporcionaram o apoio financeiro para que fosse possível eu me dedicar a esta pesquisa.

Aos demais familiares e amigos que estiveram comigo nos momentos bons e ruins, muita gratidão.

Que Deus esteja sempre conosco.

RESUMO

A variabilidade da precipitação na Amazônia Legal, especialmente em áreas com alta incidência de zeros nas séries temporais, é um problema significativo para a modelagem climática e a previsão precisa. Este estudo teve como objetivo modelar e prever a precipitação na Amazônia Legal utilizando o Modelo Aditivo Generalizado de Localização, Escala e Forma (GAMLSS), com foco na análise de teleconexões, velocidade do vento, pressão atmosférica e temperaturas extremas como variáveis explicativas. Foram coletados e analisados dados de precipitação e outras variáveis climáticas durante os trimestres de 2021. O modelo GAMLSS foi empregado para identificar padrões regionais de precipitação e a influência de teleconexões, utilizando a distribuição *Zero Adjusted Gamma* (ZAGA) para lidar com a alta incidência de zeros nas séries temporais. Os resultados indicam que a distribuição ZAGA foi particularmente eficaz, alcançando valores de R^2 superiores a 0,75 em 132 dos 408 pixels analisados. As regiões do Pará, Maranhão, leste do Amazonas e norte do Tocantins, juntamente com áreas no Acre, Rondônia e Mato Grosso, destacaram-se pelo alto desempenho do modelo. O Pará, em particular, abrigou 44% dos modelos com R^2 superiores a 0,90. A conclusão reforça a robustez da distribuição ZAGA na previsão de precipitação em regiões com alta variabilidade climática, oferecendo uma base sólida para futuras pesquisas e políticas de conservação climática na Amazônia Legal.

Palavras-chave: Teleconexões. GAMLSS. Predição. Amazônia Legal.

ABSTRACT

The variability of precipitation in the Legal Amazon, particularly in areas with a high incidence of zeros in time series, is a significant challenge for climate modeling and accurate forecasting. This study aimed to model precipitation in the Legal Amazon using the Generalized Additive Model for Location, Scale, and Shape (GAMLSS), focusing on the analysis of teleconnections, wind speed, atmospheric pressure, and extreme temperatures as explanatory variables. Precipitation and other climatic variables were collected and analyzed during the quarters of 2021. The GAMLSS model was employed to identify regional precipitation patterns and the influence of teleconnections, using the Zero Adjusted Gamma (ZAGA) distribution to handle the high incidence of zeros in the time series. The results indicate that the ZAGA distribution was particularly effective, achieving R^2 values higher than 0.75 in 132 out of the 408 pixels analyzed. The regions of Pará, Maranhão, eastern Amazonas, and northern Tocantins, along with specific areas in Acre, Rondônia, and Mato Grosso, stood out for the high performance of the model. Pará, in particular, accounted for 44% of the models with R^2 values above 0.90. The conclusion highlights the robustness of the ZAGA distribution in forecasting precipitation in regions with high climate variability, providing a solid basis for future research and conservation policies in the Legal Amazon.

KEYWORDS: Teleconnections, GAMLSS, predictions, Brazilian Legal Amazon.

SUMÁRIO

1	INTRODUÇÃO	16
2	OBJETIVOS	19
2.1	OBJETIVO GERAL	19
2.2	OBJETIVOS ESPECÍFICOS	19
3	REFERENCIAL TEÓRICO	20
3.1	PRECIPITAÇÃO	20
3.2	TELECONEXÕES	21
3.3	<i>MODELO ADITIVO GENERALIZADO DE LOCAÇÃO, ESCALA E FORMA</i> <i>(GAMLSS)</i>	24
3.3.1	<i>Modelos Lineares (Linear Model – LM)</i>	25
3.3.2	<i>Modelos Lineares Generalizado (Generalized Linear Model - GLM)</i>	26
3.3.3	<i>Modelos Aditivos Generalizados (Generalized Additive Model - GAM)</i>	27
3.3.4	<i>Modelos Aditivos Generalizados para Locação, Escala e Forma (GAMLSS)</i> 28	
3.3.4.1	Suavização no GAMLSS	30
3.3.4.2	Máxima Verossimilhança ou <i>Maximum Likelihood</i> (ML)	31
3.3.4.3	Máxima Verossimilhança Penalizada.....	32
3.3.4.4	Algoritmos de estimação de β e γ para um λ fixo	32
3.3.4.5	Estimando o hiperparâmetro λ	34
3.3.4.6	Seleção de Modelo.....	34
3.3.4.7	Análise dos Resíduos.....	38
4	MATERIAIS E MÉTODOS	44
4.1	CARACTERÍSTICAS DA ÁREA DE ESTUDO	44
4.2	ANOMALIAS DAS TELECONEXÕES	46
4.3	EXTRAÇÃO DAS DEMAIS VARIÁVEIS EXPLICATIVAS.....	47
4.4	CONSTRUÇÃO DOS MODELOS	49
4.5	MÉTRICAS DE ANÁLISE	51
4.5.1	<i>Índice Kappa</i>	52
4.5.2	<i>Índice Kendall</i>	53
5	RESULTADOS E DISCUSSÕES	55
5.1	ANÁLISE DESCRITIVA DOS DADOS DE PRECIPITAÇÃO.....	55
5.2	CÁLCULO DO R^2 EM CADA MODELO DESENVOLVIDO	56
5.3	ANÁLISE QUALITATIVA DO MODELO	72
6	CONCLUSÕES.....	80
	REFERÊNCIAS BIBLIOGRÁFICAS	84
	APÊNDICE A	93
	APÊNDICE B.....	122

LISTA DE FIGURAS

Figura 1 - Funcionamento da estimação dos parâmetros para uma distribuição Weibull (μ σ) por meio dos algoritmos (a) RS e (b) CG (Fonte: STASINOPOULOS *et al.*, 2017, p. 63) - Página 12

Figura 2 - Descrição de como o resíduo quantílico normalizado é obtido por meio de uma função contínua (Fonte: STASINOPOULOS *et al.*, 2017, p. 420) - Página 37

Figura 3 – Diferentes diagnósticos do Worm Plot (Fonte: Autor) - Página 39

Figura 4- Delimitação territorial da Amazônia Legal (Fonte: Autor) - Página 41

Figura 5- Delimitação territorial da Amazônia Legal (Fonte: Autor) - Página 42

Figura 6 - Resumo gráfico da metodologia na construção do modelo (Fonte: Autor) - Página 46

Figura 7 -Detalhes das distribuições- Página - 47

Figura 8 - Distribuição dos Pixels na Amazônia Legal Brasileira (Fonte: Autor) - Página 48

Figura 9 – R² dos modelos desenvolvidos pela família ZAGA (Fonte: Autor) - Página 53

Figura 10 – R² de valores superior a 0.75 (Fonte: Autor) - Página 54

Figura 11 – – R² de valores superior a 0.90 (Fonte: Autor) - Página 55

Figura 12 - Distribuição da precipitação nos 27 melhores pixels – Parte 1 (Fonte: Autor) - Página 57

Figura 13 - Distribuição da precipitação nos 27 melhores pixels – Parte 2 (Fonte: Autor) - Página 59

Figura 14 - Distribuição da precipitação nos 27 melhores pixels – Parte 3 (Fonte: Autor) - Página 60

Figura 15 - Correlações e Tendências das Variáveis Explicativas (Fonte: Autor) - Página 61

Figura 16 – Worm plot dos pixels 233-301 (Fonte: Autor) - Página 62

Figura 17 – Worm plot dos pixels 306-369 (Fonte: Autor) - Página 63

Figura 18 – Worm plot dos pixels 387-575 (Fonte: Autor) - Página 64

Figura 19 - Resultado da análise qualitativa para o 1º Trimestre (Fonte: Autor) - Página 70

Figura 20 - Resultado da análise qualitativa para o 2º Trimestre (Fonte: Autor) - Página 71

Figura 21 - Resultado da análise qualitativa para o 3º Trimestre (Fonte: Autor) - Página 71

Figura 22 - Resultado da análise qualitativa para o 4º Trimestre (Fonte: Autor) - Página 72

Figura 23 - Valores de Kappa (Fonte: Autor) - Página 73

Figura 24 - Valores de Kappa igual a 1 (Fonte: Autor) - Página 74

Figura 25 - Valores de Kendall (Fonte: Autor) - Página 75

Figura 26 - Valores de Kendall iguais a 1 (Fonte: Autor) - Página 76

LISTA DE TABELAS

Tabela 1 - Interpretação de distintos padrões do gráfico de minhoca - Página 40

Tabela 2 - Intervalos correspondentes às classificações climáticas - Página 50

Tabela 3 - Importância das anomalias de teleconexões nos 27 melhores modelos - Página 57

Tabela 4 - Importância da temperatura e vento como variáveis explicativas nos 27 melhores modelos - Página 58

1 INTRODUÇÃO

A precipitação é um fenômeno meteorológico resultante do processo de condensação do vapor de água, influenciado por diversos processos atmosféricos complexos, que interagem com uma variedade de fatores atmosféricos e oceânicos ocorrendo em diferentes escalas de tempo e espaço (Canchala et al., 2020). Este fenômeno desempenha um papel crucial como principal fonte de aporte de água nas bacias hidrográficas (Muhammad et al., 2018) e é uma das variáveis hidrológicas mais relevantes para a gestão de recursos hídricos (Birsan et al., 2005; Evans, Bennett e Ewenz, 2009; Manatsa, Chingombe e Matarira, 2008; Rashid et al., 2015).

A importância do estudo da precipitação estende-se a vários setores da sociedade, como a agricultura (Trinh, 2018), o setor energético — principalmente por meio das hidrelétricas (Haddad, 2011; Santos et al., 2019) — e a gestão hídrica (Hartmann et al., 2016; Serinaldi e Kilsby, 2012). Dada sua relevância, a previsão precisa da precipitação é um dos grandes desafios enfrentados pela hidrologia moderna.

As teleconexões referem-se a padrões climáticos de grande escala que conectam variações atmosféricas e oceânicas em diferentes regiões do planeta, influenciando o clima global de forma interligada. Esses fenômenos ocorrem quando mudanças em uma região específica afetam as condições climáticas de áreas distantes. No contexto da Amazônia Legal, teleconexões desempenham um papel crucial, modificando os padrões de precipitação e temperatura na região. Tais conexões influenciam diretamente a dinâmica do regime hídrico amazônico, resultando em variações na intensidade das chuvas e na ocorrência de eventos extremos, como secas prolongadas ou inundações. Dessa forma, compreender essas interações é essencial para a previsão de precipitação na Amazônia e para mitigar os impactos das mudanças climáticas na maior floresta tropical do mundo (Wang et al., 2020).

A influência dessas teleconexões, associada a outros fatores, como o aumento da industrialização, poluição e urbanização, vem se mostrando como uma variável explicativa e, em alguns estudos, se mostra como intensificador na ocorrência de desastres naturais, incluindo inundações e secas. A precipitação, diretamente afetada por essas teleconexões, está fortemente associada a desastres naturais que causam sérios danos econômicos, sociais e ambientais (Pham, 2020; Castro e Calheiros, 2007; Bawa e Seidler, 2015; Heald e

Spracklen, 2015). O aumento da frequência e severidade desses eventos é agravado pelas mudanças no uso da terra e pelas condições climáticas alteradas, o que torna ainda mais urgente o estudo e a previsão precisa da precipitação na Amazônia Legal (Bezack et al., 2016; Bui et al., 2019; Faccini et al., 2018; Janizadeh et al., 2019; Deo et al., 2018; Mouatadid et al., 2018).

As secas, por exemplo, são originadas por déficits de precipitação e representam um dos maiores desafios para a gestão de recursos hídricos, com impactos profundos em vários setores (Ionita et al., 2016). No Nordeste do Brasil, Dantas et al. (2020) demonstraram a influência das teleconexões sobre a precipitação, desenvolvendo um modelo GAMLSS capaz de prever a chuva com base nas anomalias dessas teleconexões. Estudos similares, como os de Phillips et al. (2012), Huang et al. (2016) e Ni et al. (2017), evidenciaram a relação entre teleconexões, como o El Niño-Oscilação Sul (ENSO) e a Oscilação Ártica (AO), com variações nos armazenamentos de água em diferentes regiões, confirmando a influência dessas conexões no regime hídrico global.

Estudos mais recentes também reforçam essa relação entre teleconexões e variabilidade hídrica. Liu et al. (2020) investigaram essa dinâmica em regiões da Ásia e Leste Europeu, enquanto Kumar et al. (2021) e outros pesquisadores observaram a sensibilidade dos modelos climáticos às teleconexões e sua capacidade de prever eventos extremos e mudanças nos padrões de precipitação (Gao et al., 2018; Rashid e Beecham, 2019). Essas descobertas têm implicações diretas para a modelagem estatística de séries temporais, como os modelos GAMLSS, que se mostram eficazes em capturar a influência de variáveis climáticas sobre a precipitação (Bayer e Castro, 2012; Fathian et al., 2019).

A Amazônia Legal, com sua vasta biodiversidade e papel crucial na regulação do clima regional e global, também sofre os impactos das variações climáticas. A evapotranspiração intensa da região contribui para a circulação atmosférica global (Marengo e Betts, 2011; Almeida et al., 2016), e mudanças no uso da terra, como o desmatamento, podem alterar significativamente os fluxos de energia e padrões de precipitação (Lejeune et al., 2015). Modelos climáticos sugerem que a expansão do desmatamento pode elevar as temperaturas e reduzir a precipitação, aumentando os efeitos negativos sobre o clima e os ecossistemas regionais (Marengo, 2004).

Portanto, estudar a precipitação na Amazônia Legal é essencial não apenas para a conservação da biodiversidade e dos serviços ecossistêmicos que a floresta oferece, mas também para prever e mitigar os impactos das mudanças climáticas globais nesta região crítica. Diante dessa problemática, este trabalho buscará definir um sistema de previsão capaz de antecipar ações para mitigar os efeitos da variabilidade temporal. Para isso, será utilizado o modelo aditivo generalizado de localização, escala e forma (GAMLSS), embasado nas relações entre as teleconexões atmosféricas, pressão atmosférica e temperaturas extremas (Marengo *et al.*, 2017; Bridgman e Oliver, 2014; Kumar, 2021).

2 OBJETIVOS

2.1 Objetivo Geral

Modelar e prever a precipitação na Amazônia Legal utilizando o *Modelo Aditivo Generalizado de Localização, Escala e Forma* (GAMLSS), com foco na análise de teleconexões, velocidade do vento, pressão atmosférica e temperaturas extremas como variáveis explicativas

2.2 Objetivos Específicos

- Coletar e analisar dados de precipitação para identificar padrões e variações regionais ao longo dos trimestres do ano de 2021.
- Determinar quais teleconexões têm maior influência nas respostas dos modelos preditivos de precipitação na área de interesse.
- Analisar a aplicação do GAMLSS estritamente paramétrico com a distribuição ZAGA ao conjunto de dados estudados, avaliando métodos de estimação, distribuições, seleção de modelos e análise de resíduos para cada gride analisado.

3 REFERENCIAL TEÓRICO

3.1 PRECIPITAÇÃO

A precipitação é toda água advinda do meio atmosférico que chega à superfície terrestre podendo apresentar-se como granizo, chuva ou neve (Collischonn e Dornelles, 2015; Tucci, 2001). Na hidrologia, o entendimento deste fenômeno é fundamental para o melhor gerenciamento hídrico, tendo em vista que, dentre as variáveis hidrológicas a precipitação é a de maior relevância na gestão de recursos hídricos (Birsan *et al.*, 2005; Evans; Bennett; Ewenz, 2009; Manatsa; Chingombe; Matarira, 2008, Rashid *et al.*, 2015).

Desde os primórdios das civilizações mais antigas, o homem já se preocupava com as análises climáticas. Isso é explicado pelo fato da inerente influência dos fenômenos atmosféricos com a forma de viver da sociedade. Com isso, há uma busca pelo melhor controle e entendimento dos recursos hídricos, já que estes apresentam forte ligação com a melhor administração governamental e, conseqüentemente, melhor desenvolvimento socioeconômico, mostrando-se como pilar fundamental para a sociedade contemporânea.

Sendo a precipitação a principal fonte de entrada do ciclo hidrológico e tendo destaque por sua grande importância das produções agrícolas, o entendimento e a sua caracterização em termos de variabilidade espaço temporal são imprescindíveis para uma melhor administração climática e social (Schneider *et al.*, 2016; Muhammad *et al.*, 2018, Teodoro *et al.*, 2016).

Tradicionalmente, os dados de precipitação são coletados por meio de estações pluviométricas, utilizando instrumentos como pluviômetros ou pluviógrafos. Segundo Cao (2018), essas fontes representam a principal e mais direta forma de obtenção de dados sobre precipitação. No entanto, o uso exclusivo de dados provenientes de postos pluviométricos apresenta limitações significativas.

As limitações das medições realizadas por postos pluviométricos estão relacionadas à natureza pontual dessas observações, o que dificulta o estudo e a compreensão da distribuição espacial da chuva. Além disso, a dependência de processos mecânicos e anotações manuais para registro de dados pode resultar em erros e falhas nas medições

(Pereira *et al.*, 2013; Curtarelli *et al.*, 2014; Rao *et al.*, 2015; Brasil Neto, 2020; Goovaerts, 2000; Soares, 2016; Terink *et al.*, 2018).

A manutenção e operação dos postos pluviométricos também enfrentam desafios significativos devido aos altos custos envolvidos, o que pode resultar em estações com defeitos ou fora de operação. Esses fatores, combinados com a natureza pontual e não contínua das medições, contribuem para as limitações desse método de coleta de dados.

A rede global e nacional de postos pluviométricos é espacialmente irregular, sendo muitas vezes insuficiente para um estudo abrangente e detalhado da distribuição espaço-temporal da precipitação. No Brasil, essa irregularidade é evidenciada pela presença de mais de 25 diferentes regiões pluviométricas, refletindo a diversidade dos regimes climáticos do país e ressaltando a necessidade crítica de dados confiáveis (Brasil Neto, 2020).

Frente a estas questões, ao longo dos anos, têm surgido e ganhado destaque métodos alternativos para a aquisição de dados de chuva. Entre esses métodos, a utilização de produtos de sensoriamento remoto orbital, que permitem estimativas sistemáticas da precipitação em amplas áreas geográficas. Assim, esses produtos surgem como uma valiosa fonte de dados para estudos no campo dos recursos hídricos. Essa busca por dados precisos de estimativa de precipitação tem despertado o interesse de diversos pesquisadores, impulsionando não apenas o aprimoramento das técnicas de coleta de dados, mas também o monitoramento de recursos naturais em abordagens multidisciplinares.

3.2 TELECONEXÕES

A interação entre dois dos mais importantes componentes do sistema climático da Terra: atmosfera e oceano, desempenha um papel fundamental na regulação do clima global e regional. Os oceanos absorvem energia solar e esse calor é transferido para a atmosfera por meio da evaporação. Essa transferência influencia diretamente os padrões de circulação atmosférica, como na formação de sistemas de alta e baixa pressão e na movimentação de massas de ar.

Os padrões de circulação atmosférica e oceânica referem-se aos movimentos e fluxos de ar na atmosfera e de água nos oceanos ao redor do planeta. A circulação atmosférica pode

ser expressada como o movimento horizontal e vertical do ar na atmosfera terrestre. Esse movimento é impulsionado principalmente pelo aquecimento desigual da superfície da Terra pelo sol. As principais características da circulação atmosférica incluem as células de circulação atmosférica, os ventos alísios e os padrões de circulação de grande escala. As células de circulação atmosféricas são padrões de movimento do ar que se estabelecem devido às diferenças de temperatura e pressão atmosférica entre o equador e os polos. Os ventos alísios são ventos constantes que sopram dos trópicos em direção ao equador e aos polos. Já os padrões de circulação de grande escala são sistemas de alta e baixa pressão, frentes atmosféricas e correntes de jato, que influenciam diretamente o clima diário e sazonal em várias partes do mundo.

Por sua vez, a circulação oceânica refere-se aos movimentos de água nos oceanos, que são influenciados por diferenças de temperatura, salinidade, ventos e pelo efeito de rotação da Terra. Os principais elementos da circulação oceânica incluem as correntes oceânicas, a ressurgência e o efeito termodinâmico. As correntes oceânicas são fluxos de água que se movem ao redor dos oceanos, levando nutrientes e transportando calor. A ressurgência é definida como o movimento ascendente de águas frias e ricas em nutrientes para a superfície. Por último, o efeito termodinâmico está relacionado com o papel fundamental dos oceanos de redistribuir o calor ao redor do planeta

Os padrões de circulação atmosférica e oceânica são fundamentais para o transporte de calor e umidade, influenciando diretamente o clima e o tempo em diferentes regiões da Terra. Dessa forma, as teleconexões ou fenômenos climáticos remotos, são padrões ou eventos climáticos que ocorrem em uma região do globo e possui influência significativa em outra região distante geograficamente. Essas influências podem ocorrer devido a interações atmosféricas e oceânicas em larga escala que interligam diversas partes do mundo através de alterações nos padrões de circulação atmosférica e oceânica.

As teleconexões são influenciadas por variações em grande escala nos padrões de circulação atmosférica, que por sua vez são afetadas por forçantes climáticas. Estas são definidas como fatores ou processos que perturbam o balanço energético da Terra, causando alterações no clima global. Essas forçantes podem ser naturais ou causadas pelo homem e têm o potencial de influenciar o clima de várias maneiras.

As mudanças na circulação atmosférica podem provocar alterações e efeitos climáticos, agindo sobre os ventos, chuvas e na Temperatura da Superfície do Mar (TSM) de maneira a atingir diferentes regiões do mundo, demonstrando a complexidade das interações entre os vários componentes do sistema climático global.

Por isso, as teleconexões atmosféricas e seus efeitos têm sido a base para uma série de estudos observacionais em escala global ou em áreas específicas (ZHANG *et al.*, 1997; FOLLAND *et al.*, 2009, GRIMM e SABOIA, 2015; SANTOS, 2019). Essas pesquisas reiteram que essas oscilações estão relacionadas a fenômenos nos campos oceânicos e atmosféricos (GRIMM e SABOIA, 2015).

Ao longo do tempo, os padrões atmosféricos sofrem mudanças e são moldados por várias interações entre diferentes conexões atmosféricas em escalas temporais que vão desde anos até décadas. Nesse contexto, a Oscilação Decadal do Pacífico (PDO) e a Oscilação Multidecadal do Atlântico (AMO) emergem como os principais determinantes nas escalas de décadas.

Em escalas anuais, fenômenos como a Oscilação do Atlântico Norte (NAO), o Padrão Pacífico-América do Norte (PNA), o Índice Multivariável de ENSO (MEI), o Índice de Oscilação do Sul (SOI), o Atlântico Sul Tropical (TSAI), o Atlântico Norte Tropical (TNAI) e a Oscilação Antártica (AAO) contribuem para modular a variabilidade climática em nível regional. A dinâmica da PDO e da AMO envolve interações com a temperatura da superfície do mar (TSM) tropical, fornecendo insights cruciais sobre os impactos anuais dos eventos quentes e frios do El Niño-Oscilação Sul (ENSO) conforme as fases desses padrões climáticos.

A Oscilação Multidecadal do Atlântico (AMO) é um exemplo significativo de teleconexão que modula a temperatura da superfície do mar no Atlântico Norte ao longo de décadas. Suas fases positivas e negativas têm implicações diretas nos padrões climáticos em diversas regiões, incluindo a intensificação de fenômenos como furacões.

O Índice de Oscilação do Sul (SOI) é outro indicador importante, medindo a diferença de pressão atmosférica entre Taiti e Darwin. Ele desempenha um papel crucial na monitorização do El Niño-Oscilação Sul (ENSO), um fenômeno complexo que altera as condições oceânicas e atmosféricas no Pacífico tropical, afetando o clima global.

A Oscilação do Atlântico Norte (NAO) e a Oscilação Ártica (AO) são teleconexões que influenciam os padrões de vento e clima na região do Atlântico Norte e nas latitudes elevadas, respectivamente. Suas variações têm ramificações significativas na temperatura e precipitação em áreas como Europa e América do Norte.

A Oscilação Decadal do Pacífico (PDO) é conhecida por suas mudanças de longo prazo na temperatura da superfície do mar no Oceano Pacífico, afetando a variabilidade climática em escalas decadais e influenciando padrões de precipitação em várias regiões do mundo.

O El Niño-Oscilação Sul (ENSO) é um fenômeno climático natural que ocorre periodicamente no Oceano Pacífico tropical. Ele envolve interações complexas entre a atmosfera e o oceano e está associado a variações na temperatura da superfície do mar (TSM) e nos padrões de vento.

Durante um evento de El Niño, as águas da superfície do Pacífico tropical central e oriental se aquecem significativamente, alterando os padrões atmosféricos globais. Isso influencia os ventos alísios (ventos que sopram dos trópicos em direção ao Equador) e a circulação atmosférica, causando impactos climáticos em diferentes partes do mundo.

A medição da variação da temperatura nos oceanos pode ser feitas em áreas distintas. Dessa forma, os índices NINO1+2, NINO3, NINO3.4 e NINO4 representam diferentes regiões ao longo do equador no Pacífico e são cruciais para monitorar e prever eventos de El Niño e La Niña. Esses fenômenos têm um impacto significativo nos padrões climáticos globais, afetando a temperatura e precipitação em escala mundial.

3.3 *MODELO ADITIVO GENERALIZADO DE LOCAÇÃO, ESCALA E FORMA (GAMLSS)*

O modelo aditivo generalizado de locação, escala e forma (GAMLSS) apresenta-se como uma evolução dos demais modelos de regressão mais discutidos na literatura, quais sejam: o Modelo Linear (LM), o Modelo Linear Generalizado (GLM) e o Modelo Aditivo Generalizado (GAM) (Rigby e Stasinopoulos; 2005).

Dessa forma, para melhor entendimento do modelo GAMLSS em si, é interessante discutir e entender a evolução dos modelos precedentes que servem de base para o cálculo estatístico, desde das estruturas de cálculo mais simples até as mais complexas.

3.3.1 Modelos Lineares (Linear Model – LM)

Por muitos anos, para grande parte dos fenômenos aleatórios que necessitavam ser descritos, foram utilizados os modelos lineares. No século XIX no ano de 1809, Carl Friedrich Gauss publicou um artigo, em que se foi feito um estudo das combinações das observações dos erros, culminando no princípio do mínimos quadrados. A curva encontrada nas análises dos erros ficou conhecida como curva normal e impulsionou a formulação da regressão linear clássica. A origem da análise de regressão foi idealizada pelo estatístico Francis Galton que em 1886 estudou se havia tendência entre a altura dos filhos com a altura dos pais, observou-se a altura como sendo uma variável que tendia à altura média da população e não do pai (Turkman; Silva, 2000).

O modelo linear proposto tem sido a técnica principal de modelagem estatística por mais de um século. Este tipo de modelo é utilizado para estudar as relações entre duas variáveis: uma explicativa e a outra como variável resposta. Todavia, são necessárias algumas suposições para a aplicação do modelo, como por exemplo o pressuposto de que a variável resposta deve seguir uma distribuição normal, que os erros devem ser homoscedásticos - ou seja, que não haja variação na variância dos erros analisados - e que os erros sejam normalmente distribuídos e independentes.

Até os dias atuais a análise de regressão linear é muito utilizada na resolução de inúmeros problemas. Por meio dela, torna-se possível o entendimento linear entre duas ou mais variáveis quantitativas (β), sendo uma a variável resposta (Y_i) e uma ou mais as variáveis explicativas (Equação 1). Este modelo representa a variável independente como uma soma de combinações lineares de parâmetros desconhecidos com covariáveis que independem entre si e com um erro aleatório. Além disso, o vetor aleatório referente ao erro deverá seguir distribuição normal de média zero e variância σ^2 , implicando uma distribuição da mesma família à variável resposta.

$$Y_i = \beta_0 + \beta_1 x_{i1} + \dots + \beta_r x_{ir} + \epsilon_i \quad \text{Eq.1}$$

$$\text{Onde: } \epsilon_i \stackrel{\text{ind}}{\sim} \mathcal{N}(0, \sigma^2)$$

3.3.2 Modelos Lineares Generalizado (Generalized Linear Model - GLM)

O modelo linear generalizado é um avanço do ML cuja distribuição necessariamente deveria seguir uma distribuição normal. Nelder e Wedderburn (1972) avançaram no desenvolvimento de uma proposta, na qual o comportamento entre as variáveis explanatórias e a média da variável resposta é estabelecida por meio de uma função de ligação (McCulloch; Searle, 2001). Assim, há uma mudança na relação de dependência entre a variável resposta e a variável explicativa, que deixa de ser entendida, necessariamente, de forma direta – por meio de uma ligação identidade – e passa a ser compreendida por funções denominadas de funções de ligação.

Dessa forma, o que os autores propuseram foi uma alteração na esperança da variável resposta que liga o preditor linear à parte aleatória por meio de uma função de ligação. A variável aleatória Y é relacionada com uma função de probabilidade e é considerada uma componente aleatória com esperança condicional $E(Y|X)$ igual a μ . As variáveis explicativas X e os parâmetros do modelo β formam a estrutura sistemática que é chamado de preditor linear (η). Por último, tem-se a função de ligação que irá ligar o componente aleatório ao componente sistemático, relacionando a média com o preditor linear.

$$g(\mu) = \eta = X\beta \quad \text{Eq.2}$$

A família exponencial tem um papel fundamental por permitir que sejam incorporados dados assimétricos, de natureza discreta ou contínua e duplamente limitados (Demétrio; Cordeiro, 2007). Nos modelos lineares generalizados, a variância, assimetria e

curtose são calculadas de maneira implícita por meio de sua dependência do parâmetro de locação μ , ou seja, não são obtidas explicitamente em termos das variáveis independentes (Rigby; Stasinopoulos, 2005).

Logo, o passo principal na modelagem por meio do GLM consiste na escolha desses três passos fundamentais: determinação da distribuição da variável dependente, construção da matriz modelo e escolha da função de ligação.

Em suma, o modelo apresentou três avanços quando comparado com o LM, são eles: a variável resposta poderá ter distribuição de qualquer família exponencial; o resíduo e a variável resposta, nos modelos lineares generalizados não necessitam ser normalmente distribuídas; o GLM flexibiliza a imposição de distribuição previamente instituída no modelo linear, mudando a relação entre a variável resposta e o resto da equação por meio da aplicação de uma função de ligação monotônica. Dessa forma, para encontrar os parâmetros se faz necessária a utilização de processos estatísticos de fácil implementação e de menor custo computacional.

3.3.3 Modelos Aditivos Generalizados (*Generalized Additive Model* - GAM)

Embora representem um grande avanço na flexibilização do pressuposto de normalidade apresentado no modelo linear, os modelos lineares generalizados ainda revelam limitações, tendo em vista que continuam restringindo a relação da média a uma forma linear quando relacionadas as variáveis explicativas.

Dessa forma, surgem outras alternativas por meio do uso de regressão não-paramétrica ou semi-paramétrica, que possuem mais flexibilidade quando comparadas às estruturas estritamente paramétricas (Florencio, 2010).

Dentre tais modelos, destacam-se os Modelos Aditivos Generalizados (*Generalized Additive Models* - GAM) que representam uma evolução dos modelos lineares generalizados com um preditor linear. Nesta nova abordagem, inclui-se a soma de funções de suavização não-paramétrica das covariáveis configurando a possibilidade de os próprios dados guiarem sua relação com o preditor (η), que geralmente, decorre de maneira não-linear. Os modelos

aditivos generalizados permitem solucionar questões dos resíduos assimétricos ou heterocedásticos.

O modelo aditivo generalizado (GAM) surge por meio da implementação das suavizações no GLM em 1990 por Hastie e Tibshirani. Houve um grande avanço na sua utilização depois que Wood (2006) conseguiu implementar o GAM no R com o pacote mgcv.

Tendo em vista que os erros padrões calculados só são possíveis de serem obtidos na própria covariável analisada – pressupondo-se uma linearidade – não se pode interpretar o efeito das funções de suavização para a função como um todo, por isso deve-se atentar para as interpretações do modelo. Outro ponto de destaque, é o fato de que, neste tipo de modelagem de ajustes semi-paramétricos, há uma grande robustez no desenvolvimento matemático, tornando a compreensão mais restrita e reduzindo a maioria das análises aos comportamentos gráficos dos resultados, demonstrando diversas adequações do ajuste e dos resíduos. A ideia é deixar que os dados determinem a relação entre o preditor e as variáveis explicativas, em vez de impor uma linearidade.

O modelo aditivo generalizado pode ser expresso por:

$$\mathbf{Y} \stackrel{ind}{\sim} \mathcal{FE}(\boldsymbol{\mu}, \boldsymbol{\phi}) \quad \text{Eq.3}$$

$$\boldsymbol{\eta} = g(\boldsymbol{\mu}) = \mathbf{X}\boldsymbol{\beta} + s_1(\mathbf{x}_1) + \dots + s_J(\mathbf{x}_J) \quad \text{Eq.4}$$

Onde, s_j é uma função de suavização não-paramétrica aplicada a covariável x_j para $j = 1, \dots, J$ (Stasinopoulos *et al.*, 2017). Nota-se que não são todas as covariáveis que precisam de funções de suavização, pois as mesmas têm que ser contínuas caso não sejam, a função de suavização não se aplica.

3.3.4 Modelos Aditivos Generalizados para Localização, Escala e Forma (GAMLSS)

Apesar dos inúmeros avanços alcançados com o desenvolvimento dos modelos GLM e GAM, um dos principais problemas da modelagem com distribuição de dois parâmetros - fato de nestes modelos a assimetria e a curtose serem fixadas em termos da média e da variância - não foi contornado.

Entretanto, em 2005, Rigby e Stasinopoulos propuseram uma classe de modelos de regressão ainda mais completa, denominada GAMLSS - Modelos Aditivos Generalizados para Localização, Escala e Forma. Por meio desse novo método tem-se a vantagem de modelar os parâmetros de assimetria e curtose, possibilitando uma maior flexibilidade e uma expansão ao número de distribuições consideradas, pois nesta metodologia as distribuições relacionadas com a variável explicativa tornam-se ainda mais flexíveis e abrangente não necessitando fazer parte da família exponencial.

No GAMLSS, é possível modelar não apenas a média, mas todos os parâmetros de localização, escala e forma da distribuição da variável resposta. Estes parâmetros podem ser modelados em função das variáveis explicativas. Além disso, os preditores podem inserir funções não paramétricas (como visto no GAM), efeitos aleatórios e termos aditivos, enquadrando-se como um modelo de regressão semi-paramétrica. Ou seja, o GAMLSS pode ser entendido como paramétrico, pois, requer uma suposição de distribuição paramétrica para variável resposta, mas, na modelagem dos parâmetros de distribuição o GAMLSS é encarado como semi-paramétrico, já que permite uma modelagem dos parâmetros de distribuição como uma função das variáveis explicativas, podendo se utilizar de funções de suavização não-paramétricas.

A distribuição do GAMLSS é formada por quatro parâmetros: μ é o parâmetro de localização; σ é o parâmetro de escala; ν é um parâmetro de forma da distribuição associado à assimetria; já o τ associa-se à curtose. Assim, o preditor de cada parâmetro não precisa coincidir com o restante dos preditores, logo, as matrizes X_1, X_2, X_3, X_4 poderão ser distintas.

$$Y \stackrel{ind}{\sim} \mathcal{D}(\mu, \sigma, \nu, \tau) \quad \text{Eq.5}$$

$$\begin{aligned}
\eta_1 &= g_1(\mu) = X_1\beta_1 + s_{11}(x_{11}) + \dots + s_{1J_1}(x_{1J_1}) \\
\eta_2 &= g_2(\sigma) = X_2\beta_2 + s_{21}(x_{21}) + \dots + s_{2J_2}(x_{2J_2}) \\
\eta_3 &= g_3(\nu) = X_3\beta_3 + s_{31}(x_{31}) + \dots + s_{3J_3}(x_{3J_3}) \\
\eta_4 &= g_4(\tau) = X_4\beta_4 + s_{41}(x_{41}) + \dots + s_{4J_4}(x_{4J_4}),
\end{aligned}
\tag{Eq.6}$$

A quantidade de parâmetros do GAMLSS supera a de todos os outros antecessores citados (LM, GLM, GAM). Contudo, esta quantidade de parâmetros pode ser ainda maior, basta se assumir uma distribuição de probabilidade com mais de quatro parâmetros. Para isso, é necessário a definição de modelos intermediários para a distribuição que possuam mais de um parâmetro representando o desvio s ou o u por meio de função de covariáveis.

Para cada análise de parâmetro, este poderá receber uma função de ligação diferente. No GLM, estas funções se relacionavam com a distribuição escolhida para modelar a variável explicativa. Entretanto, no GAMLSS, tal função é escolhida tendo por base os valores que o parâmetro pode assumir. Dessa forma, no GAMLSS, a ligação identidade é adotada quando o parâmetro assume valores entre $(-\infty, \infty)$, a logarítmica é escolhida para valores que variam de $(0, \infty)$, já a logística para valores de $(0, 1)$.

Uma das principais vantagens do GAMLSS é a variedade de distribuições que podem ser utilizadas, existem mais de 100 distribuições implementadas no pacote `gamlss` disponibilizado no R, sendo várias discretas, contínuas, mistas, fortemente assimétricas, platicúrticas, leptocúrticas, entre outras. Além disso, é possível criar distribuições e implementá-las no pacote.

Os quatro parâmetros de distribuição modelados no GAMLSS, podem ter seu comportamento afetado pelas variáveis explicativas de diversas maneiras, pois os parâmetros podem ser ajustados por meio de funções paramétrica lineares, não-lineares ou por funções não-paramétricas. Assim, podem ser implementados termos aditivos com P-spline, splines cúbicas, loess, polinômios fracionários, efeitos aleatórios, ajustes não-paramétricos, entre outros (Stasinopoulos *et al.*, 2017).

3.3.4.1 Suavização no GAMLSS

Para as suavizações, Stasinopoulos (2017) considera que podem ser descritas da seguinte forma:

$$s(x) = Z\gamma, \quad \text{Eq.7}$$

Onde, Z é uma matriz base para a suavização que depende da variável explicativa x ; γ é um vetor de parâmetros a serem estimados, submetidos a uma penalização denotada por: $\lambda\gamma^T G \gamma$, sendo G uma matriz conhecida $= D^T D$, em que D indica a matriz de diferenças e λ sendo um vetor de hiperparâmetros que controla o grau de suavização necessária para o ajuste.

Dessa forma, podemos generalizar a equação 6, por meio da equação 8 e 9.

$$Y | \gamma \stackrel{ind}{\sim} D(\mu, \sigma, \nu, \tau) \quad \text{Eq.8}$$

$$\begin{aligned} \eta_1 &= g_1(\mu) = X_1\beta_1 + Z_{11}\gamma_{11} + \dots + Z_{1J1}\gamma_{1J1} \\ \eta_2 &= g_2(\sigma) = X_2\beta_2 + Z_{21}\gamma_{21} + \dots + Z_{2J2}\gamma_{2J2} \\ \eta_3 &= g_3(\nu) = X_3\beta_3 + Z_{31}\gamma_{31} + \dots + Z_{3J3}\gamma_{3J3} \\ \eta_4 &= g_4(\tau) = X_4\beta_4 + Z_{41}\gamma_{41} + \dots + Z_{4J4}\gamma_{4J4} \end{aligned} \quad \text{Eq.9}$$

O modelo GAMLSS poderá ser representado com efeito aleatório como visto nas equações 8 e 9, entretanto, em caso desses efeitos não existirem, passam a um modelo GAMLSS paramétrico sem funções de suavização para efeitos aleatório, requerendo apenas a estimação dos parâmetros b como visto na equação 6. A estimação é feita por meio da máxima verossimilhança. Já quando analisado um modelo GAMLSS com efeitos aleatórios, exige-se não apenas o β , mas γ e λ .

3.3.4.2 Máxima Verossimilhança ou *Maximum Likelihood* (ML)

O processo de ajuste por meio da máxima verossimilhança ou *Maximum Likelihood* (ML) se caracteriza como uma função de estimativa que maximiza a chance do parâmetro a ser estimado ser um valor correto tendo em vista a análise da amostra estudada. Dessa forma,

dado uma amostra de valores cuja média é desconhecida, mas com família de distribuição da variável sendo estimada como a normal, por exemplo, podemos estimar valores para a média e para o desvio.

Entende-se como objetivo principal do ML a criação de diversos cenários utilizando as informações iniciais – para uma distribuição normal, por exemplo, se estimaria um desvio e uma média em cada cenário – através disso, o método testará qual desses cenários possui o valor de parâmetro mais provável, ou seja, de maior verossimilhança dado os valores da amostra analisada.

Por meio desses passos, obtêm-se os valores dos parâmetros estimados não por meio de álgebra matricial, como era realizado nos modelos lineares, mas, por meio de uma estimativa de qual valor tem a maior chance de ser o correto dado a distribuição e a respectiva criação dos diversos cenários. Aquele cenário cuja probabilidade seja a maior, será o cenário que atribuirá os valores dos parâmetros estimados. Como esses valores de probabilidade acabam sendo muito pequenos, adota-se a utilização da transformação para logaritmo.

No caso do modelo GAMLSS paramétrico, utiliza-se o *Maximum Likelihood* (ML) na estimação dos quatro parâmetro sendo expresso matematicamente pela equação 10:

$$l = \sum_{i=1}^n \ln[f(y_i|\mu_i, \sigma_i, \nu_i, \tau_i)] \quad \text{Eq.10}$$

3.3.4.3 Máxima Verossimilhança Penalizada

O modelo GAMLSS não paramétrico com efeitos aleatórios é ajustado pelo método de estimação por máxima verossimilhança penalizada, em relação à estimação de β e γ para um λ fixo. Assim, a análise matemática dar-se-á pela forma expressa na equação 11:

$$l_p = l - \frac{1}{2} \sum_{k=1}^4 \sum_{j=1}^{J_k} \gamma_{kj}^T \mathbf{G}_{kj}(\lambda_{kj}) \gamma_{kj}. \quad \text{Eq.11}$$

3.3.4.4 Algoritmos de estimação de β e γ para um λ fixo

Para o ajuste dos parâmetros β e γ com um λ fixo, maximizando a função de verossimilhança penalizada, Rigby e Stasinopoulos desenvolveram dois principais algoritmos de cálculo, o CG e o RS.

O modelo CG apresenta-se como uma generalização do algoritmo de Cole e Green criado em 1992. Este método é um algoritmo de pontuação local que é executado por meio de uma iteração externa, uma iteração interna e um algoritmo de ajuste modificado focado no ajuste de cada distribuição dos parâmetros (Rigby e Stasinopoulos, 2005).

Diferentemente do algoritmo RS, o CG precisa das derivadas das probabilidades logarítmicas com relação a cada par de parâmetros da distribuição. Assim, o algoritmo CG se utilizará das primeiras derivadas e dos valores exatos ou aproximados das derivadas segundas e derivadas cruzadas da função de verossimilhança dos dados desde que respeitem $\theta = (\mu, \sigma, \nu, \tau)$.

Todavia, para muitas funções de densidade de probabilidade explicadas por (μ, σ, ν, τ) , tais parâmetros possuem informações ortogonais, o que acarreta derivada cruzada com respeito a função de verossimilhança iguais à zero. Nestes casos, será impossível utilizar tal algoritmo e será preciso recorrer ao uso do algoritmo RS.

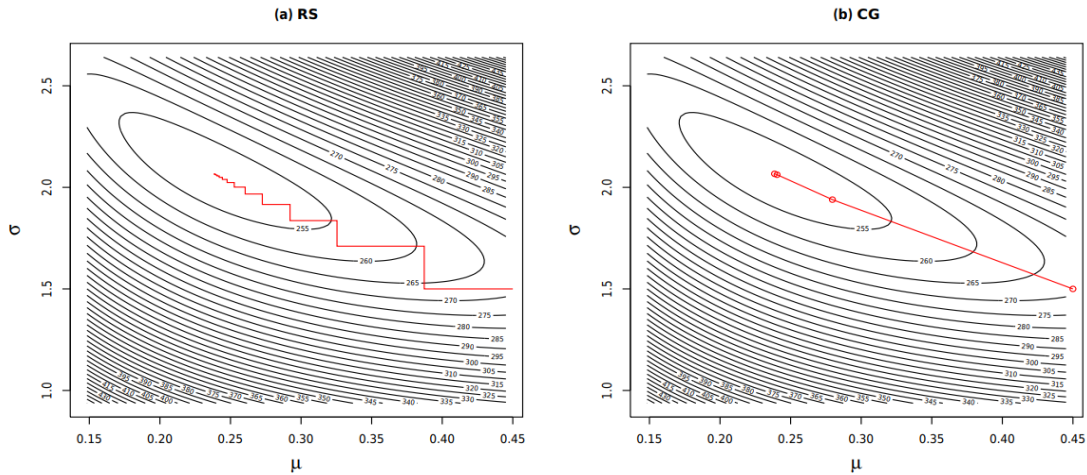
O algoritmo RS é uma generalização do algoritmo desenvolvido em 1996 por Rigby e Stasinopoulos, denominado de MADAM, utilizado para o ajuste dos modelos aditivos de média e dispersão. Neste caso, não há a utilização dos valores provenientes das derivadas. Além desses dois, existe uma terceira alternativa que é uma aplicação combinada entre os dois algoritmos, iniciando o processo iterativo com a utilização do RS e finalizando com o CG.

Quando utilizado o algoritmo RS há uma maximização da verossimilhança em cada parâmetro buscado até que se atinja uma convergência. Por outro lado, o algoritmo CG consegue atualizar todos os parâmetros de forma conjunta para cada iteração, contudo, terá a deficiência de só poder ser operacionalizado quando as derivadas cruzadas não forem nulas.

Por fim, cabe destacar a maior estabilidade, rapidez e eficiência do RS, sendo a escolha default no pacote *gamlss* do R. Em Rigby e Stasinopoulos (2005), os autores provam

que tais algoritmos levam ao máximo da função de verossimilhança penalizada, fornecendo estimadores dos parâmetros de efeitos fixos e aleatórios, $\hat{\beta}$ e $\hat{\gamma}$ para λ constante.

Figura 1 – Funcionamento da estimação dos parâmetros para uma distribuição Weibull (μ , σ) por meio dos algoritmos (a) RS e (b) CG



Fonte: (STASINOPOULOS *et al.*, 2017, p. 63)

3.3.4.5 Estimando o hiperparâmetro λ

Na seção 3.2.3.3 demonstrou-se a aplicação dos modelos RS, CG e do modelo misto para situações cujo parâmetro λ , que representa o grau de suavização, era conhecido e constante. Entretanto, existem algumas formas de se buscar estimar a suavização. Os métodos propostos iniciais eram chamados de globais porque eram aplicados fora das iterações dos algoritmos RS ou CG, e se utilizavam de otimizações numéricas para minimizar o critério de Akaike, selecionando as melhores estimativas para o hiperparâmetro λ . Contudo, isso exigia um alto custo computacional.

Com isso, Rigby e Stasinopoulos desenvolveram uma maneira de selecionar este parâmetro de forma automática, que chamam de estimação interna, que funciona dentro dos algoritmos RS e CG atuando por meio de ideias e teorias de máxima verossimilhança. Esta alternativa provou-se mais rápida em comparação às demais, e, Rigby e Stasinopoulos (2017) mostram que o desempenho é bem semelhante às técnicas anteriores de grande custo computacional.

3.3.4.6 Seleção de Modelo

Um dos principais passos na modelagem por meio do GAMLSS é a escolha da melhor distribuição para a variável resposta, dos preditores adequados para serem analisados como parâmetros da distribuição escolhida, das funções de ligações e dos hiperparâmetros.

A avaliação do modelo estatístico é relacionada com a capacidade explicativa do modelo ou preditiva que costuma ser analisada por meio da comparação com os valores dos dados teste – conjunto de dados independentes. É importante que a análise dos resultados não sejam nem sobreajustados (*overfitting*), nem subajustados (*underfitting*). De maneira geral, o que se busca é um modelo capaz de apresentar generalização, isto é, um modelo sem exagerada robustez e complexidades matemáticas, pois estes tendem a ser mais específicos, entretanto sem apresentar também grande simplicidade, visto que, se assim for, poderá implicar distanciamento da representação fidedigna do fenômeno.

Dessa forma, buscando um parâmetro estimador do melhor ajuste, desenvolve-se o desvio global GDEV, muito utilizado no GAMLSS e representado matematicamente pela equação 12.

$$GDEV = -2l(\hat{\theta}), \quad \text{Eq.12}$$

Onde $l(\theta)$ é o logaritmo da função de verossimilhança ajustada, apresentado nas Equações 10 e 11. O valor de GDEV é utilizado posteriormente no cálculo do critério de Akaike generalizado (GAIC), apresentada em Voudouris *et al.* (2012) (Equação 13)

$$GAIC(\kappa) = GDEV + (\kappa \times df), \quad \text{Eq.13}$$

Percebe-se que para cada grau de liberdade, a complexidade do modelo aumenta, assim, sendo df o valor de graus de liberdade do modelo, k é a penalidade relativa à cada grau de liberdade utilizado. Se $k = 2$ então o critério coincide com o critério de Akaike (AIC) (Akaike, 1998). Quando o $k = \ln(n)$, o critério coincide com o critério de informação bayesiano (BIC) (Schwarz *et al.*, 1978).

A ideia principal do GAIC (k) é penalizar modelos com muitos parâmetros, tendo em vista que isto poderia levar a uma construção que resultaria em *overfitting*. Assim, para algum k escolhido, quanto menor o valor de GAIC (k), melhor ajustado é considerado o modelo (Stasinopoulos *et al.*, 2017).

Vale destacar que o desenvolvimento de modelos não implica na descrição real do fenômeno. O produto matemático desenvolvido expressará uma aproximação da realidade. Assim, este não será a verdade absoluta e, conseqüentemente, não existirá um único modelo real totalmente correto. A escolha do modelo, portanto, deverá ser por meio das questões substanciais analisadas para o devido fim do estudo. Logo, cada problema terá uma estratégia distinta para a respectiva escolha do modelo ideal, cabe ao desenvolvedor do modelo escolher as melhores estratégias para criação dos melhores resultados para o determinado fim analisado (Stasinopoulos *et al.*, 2017).

Segundo Stasinopoulos (2017), as melhores estratégias de análise que buscam os melhores parâmetros, hiperparâmetros, as funções de distribuição, as funções de ligação, e as escolhas na construção dos modelos, são aquelas que possuem uma base de escolha objetiva. Mediante a isto, existem dois estágios cruciais no desenvolvimento do modelo que auxiliará as comparações e tomadas de decisão, quais sejam: o processo de ajuste e o de diagnóstico.

O processo de ajuste é a etapa em que está se analisando e comparando o desenvolvimento de diversos modelos ajustados por diversas famílias de distribuições. Nesta etapa, utiliza-se o parâmetro GAIC(k) como avaliador do melhor modelo. Na etapa de diagnóstico, tem-se a apresentação das análises do comportamento dos resíduos por meio de gráficos de dispersão de resíduos e worm plot (gráfico de minhoca) (Buuren; Fredriks, 2001) que auxiliará na detecção da inadequação ou não do modelo desenvolvido.

Por meio do pacote *gamlss* do R, existem as funções `fitDist()` e `histDist()` que auxiliam na escolha da distribuição da variável resposta. A primeira utiliza a função `gamlss()` para ajustar diferentes distribuições à variável dependente. Os argumentos da função `fitDist` são o vetor dos valores da variável dependente, o valor da penalização do critério GAIC e o tipo de distribuições a se ajustar (`'realline'`, `'realplus'` ou `'realAll'`). A função `histDist` permite visualizar diferentes distribuições ajustadas à variável dependente. Com ela, torna-

se possível a obtenção dos valores constantes para os parâmetros da distribuição, necessitando apenas a variável Y e a distribuição que deseja visualizar.

Na escolha das funções de ligações, deve-se selecionar a que garanta que os parâmetros fiquem sempre dentro os respectivos intervalos de cada uma, sendo a ligação identidade adotada quando o parâmetro assume valores entre $(-\infty, \infty)$, a logarítmica é escolhida para valores que variam de $(0, \infty)$ e a logística para valores de $(0, 1)$.

Já a seleção de variáveis explicativas é um dos assuntos mais importantes no ajuste do modelo estatístico. Para um conjunto de variáveis explicativas para consideração na modelagem do parâmetro θ_k de um modelo GAMLSS, em que $\theta = (\theta_1, \theta_2, \theta_3, \theta_4) = (\mu, \sigma, \nu, \tau)$, os parâmetros conterão fatores e variáveis que podem entrar no modelo de forma linear ou por meio de funções de suavização (Stasinopoulos *et al.*, 2017).

Existem inúmeras estratégias que visam aplicar um determinado formato de seleção das variáveis explicativas usadas na modelagem dos parâmetros do modelo. No gamlss podem ser destacados dois métodos principais de seleção dos parâmetros que são denominados de estratégias A e B. Para sua aplicação é necessário chamar a função `stepGAICAll.A()` ou `stepGAICAll.B()`.

A função mais utilizada é a `stepGAICAll.A()`. Nela, primeiro se ajusta um modelo para o parâmetro μ por meio do procedimento forward para algum GAIC selecionado, considerando os demais parâmetros constantes. O passo seguinte é ajustar σ , utilizando o mesmo procedimento, considerando o primeiro modelo para μ constante. Se houverem mais parâmetros, então ν é ajustado, considerando os modelos de μ e σ . Este procedimento ocorre até que se modele todos os parâmetros necessários.

A seleção de modelos para o último parâmetro da distribuição ocorre uma única vez. Dessa forma, após selecionados os parâmetros, o algoritmo faz o caminho inverso, agindo por meio de um retrocesso, reanalisando os parâmetros por meio de um backward process, reajustando os demais dados até que se retorne o reajuste do μ (Nakamura *et al.*, 2017). Ao fim deste processo, o algoritmo para e computa os valores de cada parâmetro que compõem o modelo GAMLSS final.

3.3.4.7 Análise dos Resíduos

Uma forma de analisar os resultados provenientes dos modelos construídos é calcular a diferença entre os valores obtidos que são estimados, com os valores reais observados, assim, tem-se os resíduos ordinários, expresso na equação 14. Este cálculo apresenta-se como um dos mais simples. Dele, partem algumas outras análises mais robustas como o resíduo de Pearson, o resíduo de Pearson padronizado, resíduo componente da deviance, entre outros.

De maneira geral, um resíduo é alguma medida de afastamento entre a observação e o seu valor ajustado pelo modelo, sendo esta medida usualmente escolhida com intuito de estabilizar a variância da distribuição amostral analisada ou induzir uma simetria, visando garantir uma comparabilidade dos resíduos e detectar os valores discrepantes.

$$\hat{\epsilon}_i = y_i - \hat{y}_i , \quad \text{Eq.14}$$

Apesar de prático, simples e efetivo, os resíduos ordinários são limitados, pois, se restringem às situações em que os resíduos tenham um comportamento normalizado. Entretanto, diferentemente do que acontece em modelos lineares com erros normais, nem sempre as variáveis respostas possuirão uma distribuição normal, nestes casos a distribuição dos resíduos também não terá aproximação à distribuição normal, mesmo que o modelo tenha um bom ajuste dos dados.

A falta de normalidade dos resíduos é notada principalmente quando analisado modelagem com dados discretos ou que possuam valores pequenos. Mediante a isto, Dunn e Smyth (1996) propuseram os resíduos quantílicos normalizados que é o método mais utilizados nos diagnósticos de modelos GAMLSS. Os resíduos quantílicos aleatorizados apresentam uma distribuição Normal independentemente da distribuição da variável resposta. A aleatorização baseia-se no teorema da inversa da função distribuição acumulada. A principal vantagem dos resíduos quantílicos normalizados é que, para qualquer distribuição da variável resposta, os resíduos verdadeiros sempre terão uma distribuição normal padrão quando o modelo assumido está correto. Com isso, os resíduos quantílicos

normalizados constroem uma maneira objetiva de verificação da adequação dos modelos sem que haja a necessidade de especificação da distribuição da variável independente.

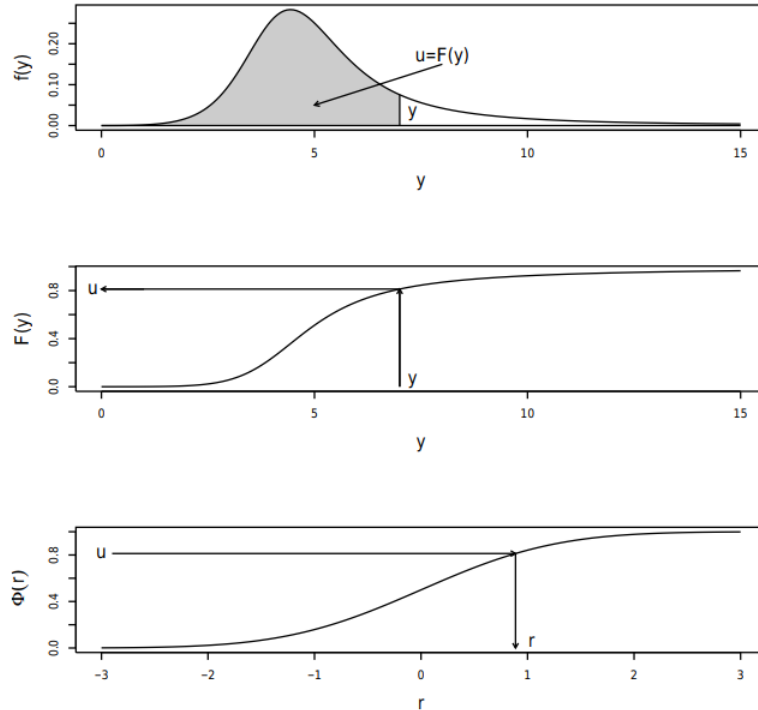
Os resíduos quantílicos normalizados podem ser representados matematicamente pela equação 15. (Stasinopoulos *et al.*, 2017).

$$\hat{r}_i = \Phi^{-1}(\hat{u}_i), \quad \text{Eq.15}$$

Onde Φ^{-1} é o inverso da distribuição acumulada de uma normal padrão e $\hat{u}_i$ são os resíduos quantílicos que se definem de maneiras distintas para variáveis discretas e contínuas.

Se y é uma observação de uma variável contínua então $u = F(y|\theta)$ e $u' = F(y|\theta')$ são os valores da função na distribuição acumulada do modelo e do ajuste, respectivamente. Se o modelo estiver bem especificado, então u terá distribuição uniforme entre os valores de zero e um. Este processo é explicado pelo teorema da transformação da integral probabilidade e é apresentado graficamente na Figura 2. O primeiro gráfico apresenta a distribuição de densidade de probabilidade para uma observação específica y . O segundo gráfico mostra como y é encontrada em u por meio da distribuição acumulada. No terceiro gráfico u é transformado em resíduo, chamado de z-score, por meio da equação 16 que terá distribuição normal padrão em caso do modelo estar correto. Da mesma maneira, u' é transformado em resíduos ajustados por r' , que possui uma distribuição aproximadamente normal, através da equação 17.

Figura 2 – Descrição de como o resíduo quantílico normalizado é obtido por meio de uma função contínua.



Fonte: (STASINOPOULOS *et al.*, 2017, p. 420)

$$r = \Phi^{-1}(u), \quad \text{Eq.16}$$

$$\hat{r} = \Phi^{-1}(\hat{u}) = \Phi^{-1} \left[F(y|\hat{\theta}) \right], \quad \text{Eq.17}$$

3.3.4.7.1 Worm Plot (Gráfico de minhoca)

O worm plot, ou gráfico de minhoca, nome dado devido ao formato do gráfico que possui, geralmente, formato de minhoca, é uma das principais técnicas utilizadas para analisar os resíduos na modelagem com o GAMLSS. Este método foi desenvolvido por Buuren e Fredriks (2001). Por meio desta ferramenta, permite-se a realização de uma checagem dos resíduos em diferentes intervalos com uma ou duas variáveis explicativas. Pode-se interpretar o gráfico como um Q-Q plot (gráfico quartil-quartil) sem tendência.

A ideia principal do worm plot é identificar regiões de uma variável explicativa dentro da qual o modelo se ajusta ou não aos dados. Os pontos plotados no gráfico mostram a distância entre os resíduos ordenados e os valores esperados (representados por meio da

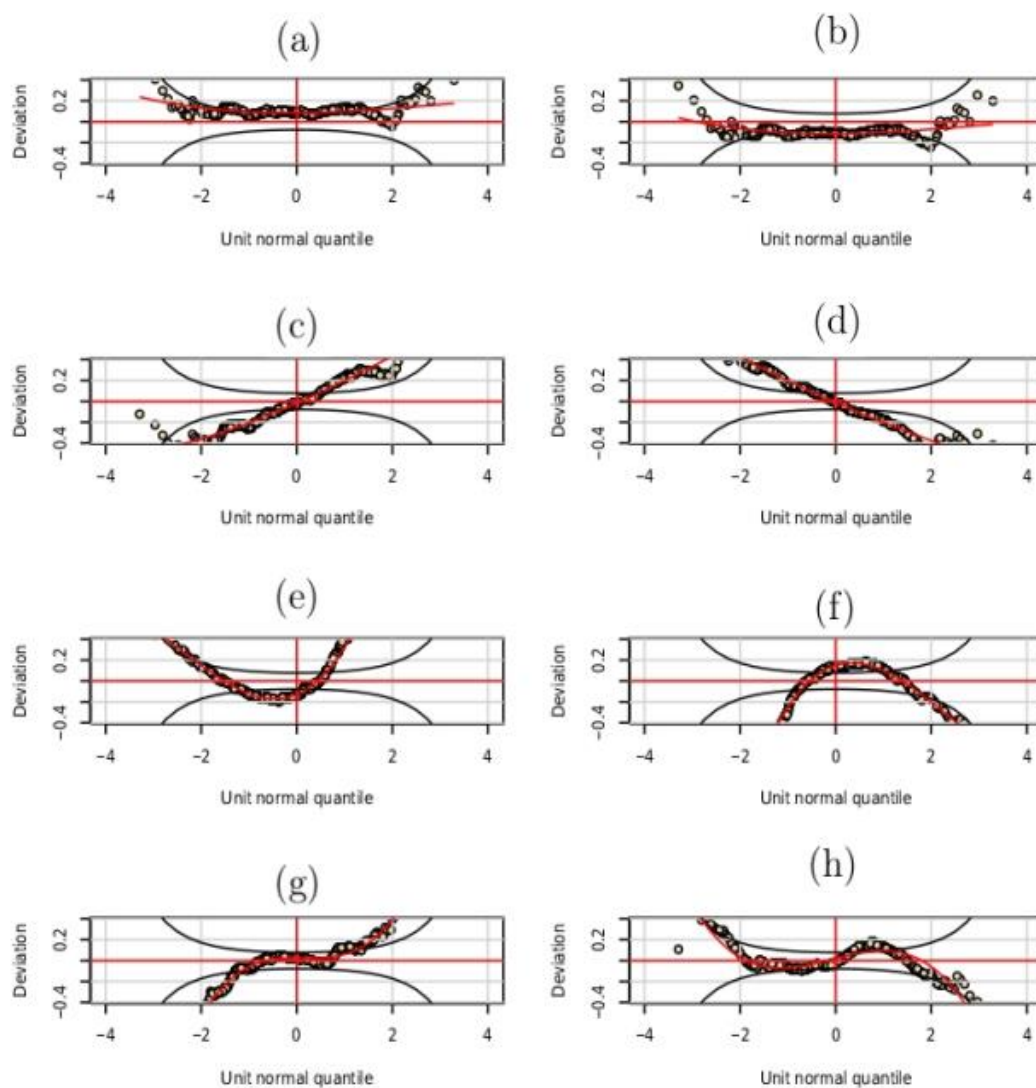
linha pontilhada horizontal). Quanto mais próximo os pontos estão da linha horizontal, mais próxima a distribuição dos resíduos estão de uma distribuição Normal padrão.

Os intervalos de confiança pontuais são de aproximadamente 95% e são dados pelas representações das curvas em elipses. Assim, para um modelo correto, teremos no máximo 5% dos pontos fora das elipses. Em caso de grande distribuição dos pontos fora das elipses, tem-se que o modelo se apresenta inadequado para explicar a variável resposta. Além disso, há a indicação de outliers, quando estes pontos estiverem muito distantes das curvas elípticas.

A curva ajustada aos pontos no worm plot é um ajuste cúbico. A forma do ajuste reflete as inadequações do modelo. Assim, em caso do modelo apresentar níveis dos pontos de plotagem no worm plot acima da linha horizontal da origem, isso indicará que o a média residual estará alta, o que implica que a localização da distribuição ajustada é baixa, ou seja, o parâmetro de média foi subestimado, ao ponto que os resíduos estão sempre altos (isso pode ser corrigido aumentando o parâmetro μ , se for um parâmetro de localização, ou melhorando o modelo para μ - por exemplo, tornando-o mais flexível- ou alterando a distribuição do modelo).

Além disso, a tendência linear, a tendência de U ou U inverso ou o formato em S indicam problema de variância, assimetria ou curtose nos resíduos. Isso aponta um problema na distribuição ajustada dos parâmetros. Demais interpretações foram resumidas e apresentadas na tabela 1.

Figura 3 – Diferentes diagnósticos do Worm Plot



Fonte: (STASINOPOULOS *et al.*, 2017, p. 429)

Tabela 1 – Interpretação de distintos padrões do gráfico de minhoca

Figura	Parâmetro Analisado	Se	Então
a	Média	a minhoca passa acima da origem,	a média ajustada é muito pequena.
b		a minhoca passa abaixo da origem	a média ajustada é muito grande.
c	Variância	a minhoca tem uma inclinação positiva	a variância ajustada é muito pequena.
d		a minhoca tem uma inclinação negativa	a variância ajustada é muito grande.
e	Assimetria	a minhoca tem formato de U	distribuição ajustada é assimétrica à esquerda.
f		a minhoca tem formato de U invertido	distribuição ajustada é assimétrica à direita.
g	Curtose	a minhoca tem uma forma em S à esquerda curvada para baixo	as caudas da distribuição ajustada são muito leves.
h		a minhoca tem uma forma em S à esquerda curvada para cima	as caudas da distribuição ajustada são muito pesadas.

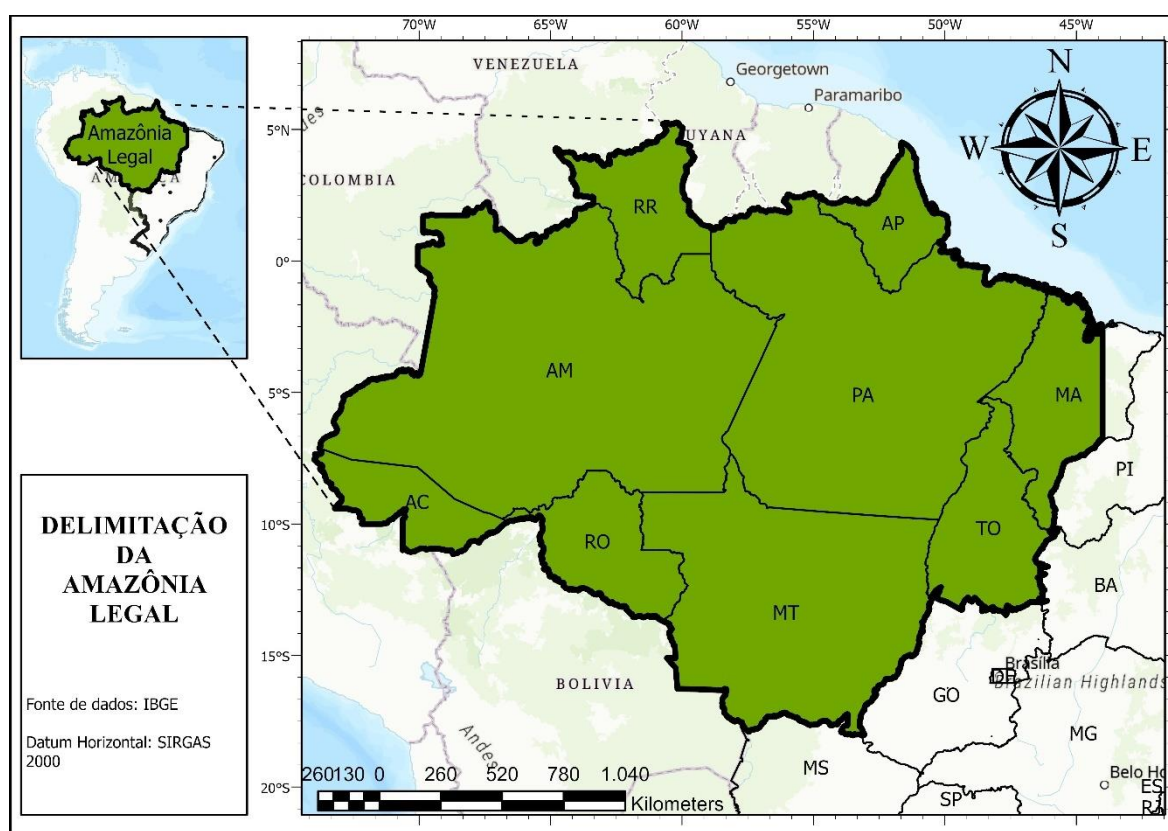
Fonte: Alcântara (2018)

4 MATERIAIS E MÉTODOS

4.1 CARACTERÍSTICAS DA ÁREA DE ESTUDO

A Amazônia Legal é uma extensa área definida pelo governo brasileiro com o propósito de promover políticas voltadas ao desenvolvimento sustentável e à preservação ambiental na região amazônica do Brasil, nela encontram-se nove estados brasileiros, quais sejam: Acre, Amapá, Amazonas, Mato Grosso, Pará, Rondônia, Roraima, Tocantins e Maranhão (Figura 4). Esta área abrange diversos biomas e ecossistemas, sendo o principal deles a floresta amazônica, conhecida por sua excepcional biodiversidade, com uma ampla variedade de espécies vegetais e animais.

Figura 4- Delimitação territorial da Amazônia Legal



Fonte: Autor

Além do bioma Amazônia, destaca-se o Cerrado, caracterizado por savanas e vegetação arbustiva e o Pantanal (7% do Mato Grosso), um dos maiores sistemas de zonas úmidas do mundo, com uma rica diversidade de vida selvagem.

O relevo da Amazônia Legal é diversificado, incluindo planícies aluviais ao longo dos grandes rios, como o Rio Amazonas e o Rio Negro, além de áreas montanhosas ao sul e sudeste, como a região da Serra do Cachimbo.

Quanto ao clima, a Amazônia Legal possui um clima predominantemente equatorial, caracterizado por temperaturas elevadas e alta umidade ao longo do ano. As estações do ano não são tão distintas como em outras regiões do Brasil, com chuvas frequentes e intensas durante a maior parte do ano.

A Amazônia Legal desempenha um papel fundamental na regulação climática global, na conservação da biodiversidade e na prestação de serviços ecossistêmicos essenciais para a vida no planeta. No entanto, enfrenta desafios significativos relacionados ao desmatamento ilegal, degradação ambiental, expansão agrícola e pressões econômicas que ameaçam sua integridade ecológica e sustentabilidade a longo prazo.

A conservação e o desenvolvimento sustentável da Amazônia Legal são temas de grande importância para o Brasil e para o mundo, exigindo a implementação de políticas eficazes e ações coordenadas para proteger e preservar esse patrimônio natural de valor inestimável.

4.2 ANOMALIAS DAS TELECONEXÕES

No presente trabalho, analisou-se os valores das anomalias das teleconexões medidas pelo *National Oceanic and Atmospheric Administration* (NOAA), que atende à missão de “compreender e prever mudanças no clima, oceanos e costas, além de buscar compartilhar esse conhecimento e informações com outras pessoas, visando gerenciar ecossistemas e recursos costeiros e marinhos”.

As atividades da NOAA vão desde previsões meteorológicas diárias, alertas de tempestades severas e monitoramento climático até gerenciamento de pesca, restauração costeira e apoio ao comércio marítimo. Os produtos e serviços da NOAA apoiam a vitalidade econômica e afetam mais de um terço do produto interno bruto dos Estados Unidos. Os cientistas da NOAA usam pesquisa de ponta e instrumentação de alta tecnologia para fornecer aos cidadãos, planejadores, gerentes de emergência e outros tomadores de decisão as informações confiáveis de que precisam.

Neste trabalho, utilizou-se a internalização das bases do NOAA, a fim de se utilizar como variável explicativa as anomalias das seguintes teleconexões: Niño 1+2, Niño 3, Niño 3.4, Niño 4, SOI, NAO, AMO e PDO.

O Niño 1+2 TSM (Temperatura de Superfície do Mar) do Pacífico Tropical do Extremo Leste (0-10S, 90W-80W) que é a menor região e a mais oriental das regiões de Niño SST e corresponde à região da costa da América do Sul, onde o El Niño foi reconhecido pela primeira vez pelas populações locais (TRENBERTH, 2017).

O Niño 3 TSM do Pacífico Tropical Leste (5N-5S, 150W-90W), essa região já foi o foco principal para monitorar e prever o El Niño, mas os pesquisadores mais tarde descobriram que a região chave para as interações entre oceano e atmosfera para o ENSO fica mais a oeste (TRENBERTH, 1997).

O Niño 3.4 TSM do Pacífico Tropical Central Leste (5N-5S) (170-120W) As anomalias do Niño 3.4 podem ser consideradas como representando as TSMs equatoriais médias em todo o Pacífico, desde a linha do tempo até a costa da América do Sul (TRENBERTH, 2017). Já o Niño 4 TSM do Pacífico Tropical Leste (5N-5S, 150W-90W) O índice Niño 4 captura anomalias de TSM no Pacífico equatorial central (TRENBERTH, 2017).

O SOI é o Índice de Oscilação Sul Índice padronizado baseado na diferença de pressão normalizada entre o Tahiti e Darwin (NCDC, 2019).

O NAO representa a Oscilação do Atlântico Norte Padrão da pressão ao nível do mar caracterizada pela diferença de pressão registradas na região da Islândia e as registradas na região dos Açores.

O AMO, é a Oscilação Multidecadal do Atlântico Modo coerente de variabilidade natural que ocorre no Oceano Atlântico Norte com um período estimado de 60–80 anos (TRENBERTH, 2017).

O PDO é a Oscilação Decadal do Pacífico Anomalias mensais da TSM no Oceano Pacífico Norte. Já a Oscilação Antártica é representada pelo AAO com o modo de variabilidade atmosférica de baixa frequência do Hemisfério Sul.

Os valores das anomalias são disponibilizados via site oficial da agência NOAA (<https://www.cpc.ncep.noaa.gov/>). Dessa forma, definiu-se o intervalo de 1958 a 2021, pois, neste intervalo há a disponibilidade dos dados para todas as anomalias utilizadas. Os arquivos em formato csv ou xlsxs foram baixados e trabalhados.

4.3 EXTRAÇÃO DAS DEMAIS VARIÁVEIS EXPLICATIVAS

Além das anomalias das teleconexões, também foram utilizadas outras variáveis explicativas nas construções dos modelos, quais sejam: temperatura máxima, temperatura mínima e velocidade máxima dos ventos.

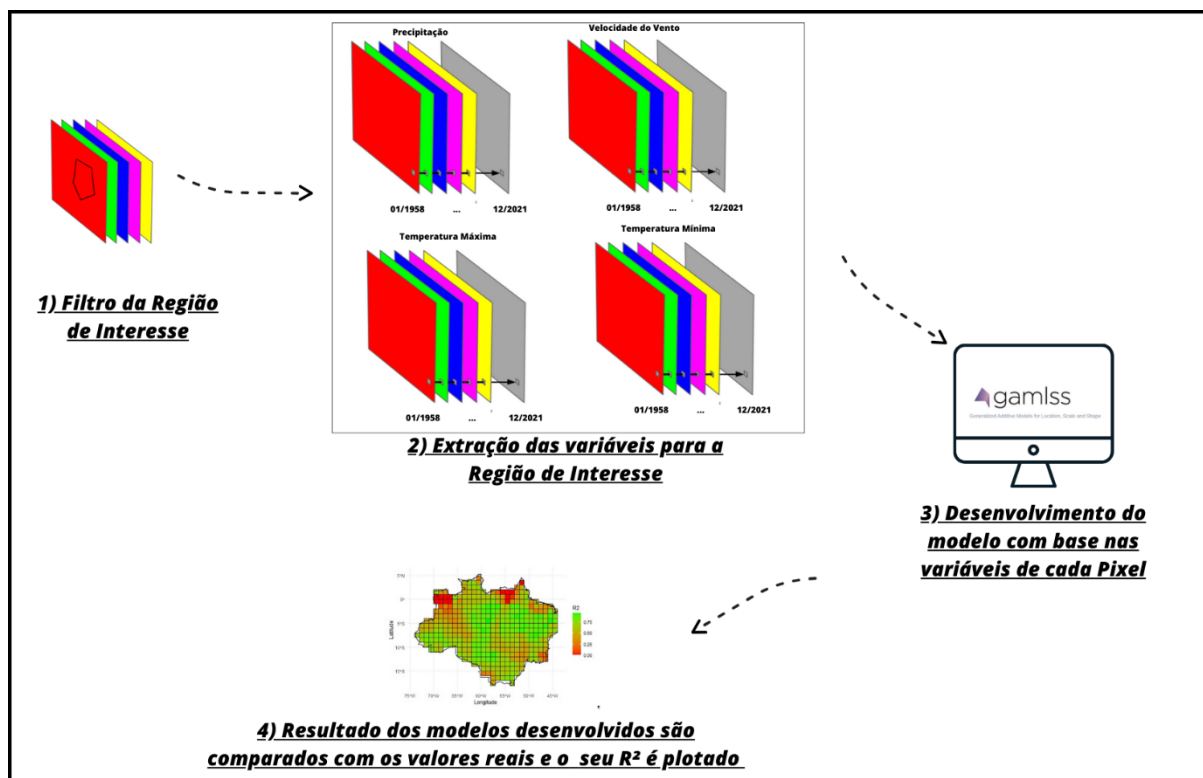
Estas variáveis são extraídas do produto TerraClimate: *Monthly Climate and Climatic Water Balance for Global Terrestrial Surfaces*, University of Idaho com resolução de 4638,3 metros. A extração dessas informações é feita por meio da utilização da IDE do *Google Earth Engine* (GEE) diante das seguintes etapas:

- I) Filtro da área de interesse e utilização das bandas (“pr” – Precipitação acumulada(mm); “tmmn” – Temperatura mínima (°C); “tmmx” – Temperatura Máxima (°C); “vs” – Velocidade do vento a 10 metros (m/s);
- II) Para cada pixel aplica-se a soma da quantidade de chuva no mês; a média da velocidade do vento a 10 metros em metros por segundo; além disso, mede-

se o valor de maior temperatura e o de menor temperatura em cada pixel presente na região de interesse;

- III) Os valores são exportados e os resultados, antes apresentados em formato de raster, podem ser representados em tabelas;
- IV) Os valores são utilizados para construção de um modelo que explicará o comportamento da precipitação para um pixel de resolução de 4638,3 metros;
- V) A imagem Geotiff contendo os valores de precipitação, temperatura máxima, temperatura mínima e velocidade dos ventos do último ano estudado (01/01/2021–01/12/2021) são retiradas da coleção, pois, esses dados serão utilizados como análise comparativa futura do comportamento dos modelos por meio da métrica R^2 .

Figura 5 – Resumo gráfico da metodologia na construção do modelo



Fonte: Autor

4.4 CONSTRUÇÃO DOS MODELOS

O processo de desenvolvimento do modelo GAMLSS passa por algumas etapas principais:

- I) O primeiro passo refere-se a escolha da família de distribuição que melhor se adequa aos dados das variáveis. O pacote gamlss apresenta uma diversidade de mais de 100 distribuições já integradas dentre elas, distribuições contínuas, discretas e mistas. Neste trabalho optou-se por utilizar a família ZAGA.
- II) O segundo passo é a escolha das funções de ligação. Estas estão relacionadas com os valores que os parâmetros podem assumir. Tendo em vista que as anomalias estão presentes nos Reais, espera-se a utilização da função identidade ou logarítmica. Com a ZAGA sendo a distribuição escolhida, tem-se que a função de ligação segue o descrito na Figura 6.

Figura 6 –Detalhe das distribuições

Distribution	gamlss name	Range R_Y	Parameter link functions			
			μ	σ	ν	τ
beta inflated (at 0)	BEOI	$[0, 1)$	logit	log	logit	-
beta inflated (at 0)	BEINFO	$[0, 1)$	logit	logit	log	-
beta inflated (at 1)	BEZI	$(0, 1]$	logit	log	logit	-
beta inflated (at 1)	BEINF1	$(0, 1]$	logit	logit	log	-
beta inflated (at 0 and 1)	BEINF	$[0, 1]$	logit	logit	log	log
zero adjusted GA	ZAGA	$[0, \infty)$	log	log	logit	-
zero adjusted IG	ZAIG	$[0, \infty)$	log	log	logit	-

Fonte: STASINOPOULOS et al., 2017

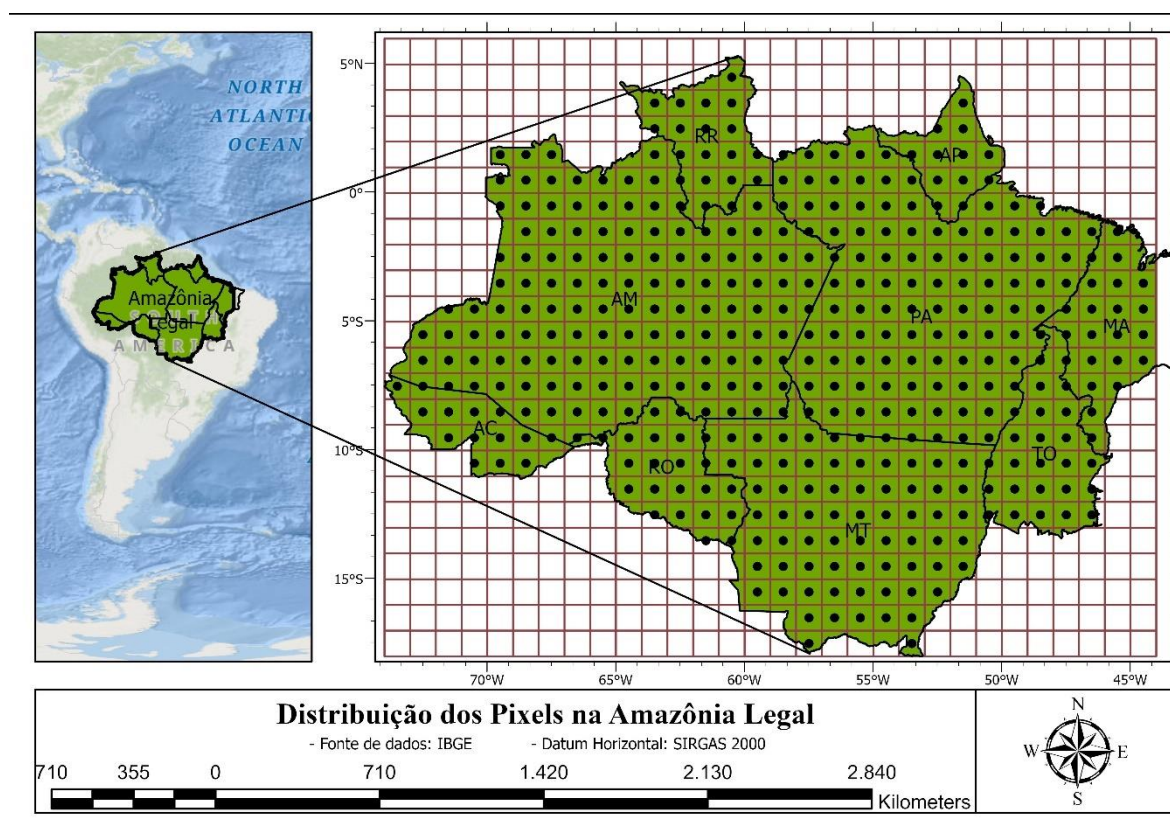
- III) Dando sequência, deve-se selecionar quais as covariáveis a serem utilizadas na formulação do modelo. Para isso, deve-se atentar para as relações previamente abordadas na literatura. Neste trabalho, optou-se em utilizar a função “stepGaicAll.A” que funciona por meio de uma estimativa de constante qualquer para os três parâmetros de escala e forma, e assume-se a inclusão de variáveis explicativas no modelo para buscar achar o melhor parâmetro de locação, avaliando isto por meio do parâmetro AIC. Ao chegar no melhor modelo de locação, trava-se o modelo e parte para a busca do parâmetro seguinte, assumindo todos os outros com valores sem variação. O

mesmo procedimento é realizado até que todos os outros parâmetros sejam calculados.

A extração dos valores de precipitação, da temperatura máxima, temperatura mínima e da velocidade do vento é feita por meio de imagem de satélite. Posteriormente, todos esses valores são convertidos em arquivos csv e unificados com os valores das anomalias das teleconexões extraídos pelo site do NOAA.

Os dados que possuem variabilidade espacial são obtidos por gride de resolução de 4638.3 metros. Analisando apenas a área de interesse, tem-se 408 grides representados na Figura 7. É desenvolvido um modelo para cada pixel e em cada um existem dados das precipitações e de todas as variáveis analisadas com resolução temporal mensal do período de 01/01/1958 até 01/12/2021.

Figura 7 – Distribuição dos pixels na Amazônia Legal Brasileira



Fonte: Autor

Tendo a família de distribuição definida sendo a ZAGA, para cada um dos 408 grides é retirado os valores do ano de 2021 para ser utilizado como dataset teste, aplica-se a função “stepGAICAll.A” a fim de buscar o modelo com as melhores variáveis explicativas para cada parâmetro.

Os valores preditos são acumulados e somados para cada trimestre do ano. Esses valores de trimestres são comparados com os valores reais e com isso é calculado a métrica do R^2 . Por fim, para uma análise espacial, cada valor de R^2 por pixel é plotado na região de interesse.

4.5 MÉTRICAS DE ANÁLISE

As métricas de análise abordadas nesse estudo avaliam a variabilidade da precipitação na área de estudo, que é influenciada por anos com precipitação significativamente acima ou abaixo da média devido à dependência das anomalias da Temperatura da Superfície do Mar (TSM). Para lidar com essa variabilidade, optou-se pelo uso da técnica dos quantis desenvolvida em R.

A aplicação da separação por quantis permite não apenas explorar se a resposta do modelo corresponde à classificação dos valores observados, mas também avaliar se o modelo consegue captar o comportamento geral da precipitação em relação à normalidade.

Esta técnica dos quantis foi inicialmente discutida por Pinkayan (1966) e subsequentemente abordada por diversos autores (Xavier *et al.*, 2002; Kayano *et al.*, 2016; Tavares *et al.*, 2018, Dantas, 2020), que demonstraram sua aplicabilidade em séries pluviométricas. Nesse contexto, a precipitação Y de uma localidade ao longo do tempo é tratada como uma variável contínua, possibilitando a geração de classificações climáticas baseadas nas categorias dos quantis. Este método é eficaz na identificação de anomalias dos eventos climáticos de interesse para períodos específicos (dos Santos, 2018).

Para este estudo, a técnica dos quantis foi adaptada para as localidades analisadas, com o objetivo de classificar os prognósticos de precipitação em sete categorias distintas: Extremamente Seco (ES), Muito Seco (MS), Seco (S), Normalidade (N), Chuvoso (C), Muito Chuvoso (MC) e Extremamente Chuvoso (EC). A Tabela 2 descreve os intervalos correspondentes às classificações climáticas representados pelos seus quantis. Dessa forma,

torna-se possível uma análise gráfica. Além disso, a análise qualitativa possibilita analisar os valores por meio de mais dois índices: O índice Kappa e o índice Kendall

Tabela 2. Intervalos correspondentes às classificações climáticas

$P < Q(0,05)$	Abaixo dos 5%	Extremamente Seco
$Q(0,05) < P < Q(0,15)$	Entre 5% e 15%	Muito Seco
$Q(0,15) < P < Q(0,35)$	Entre 15% e 35%	Seco
$Q(0,35) < P < Q(0,65)$	Entre 35% e 65%	Normalidade
$Q(0,65) < P < Q(0,85)$	Entre 65% e 85%	Chuvoso
$Q(0,85) < P < Q(0,95)$	Entre 85% e 95%	Muito Chuvoso
$Q(0,95) < P$	Superior a 95%	Extremamente Chuvoso

Fonte: Autor

Os índices Kappa e Kendall são amplamente utilizados para analisar a variabilidade entre variáveis discretas, com o intuito de determinar se os avaliadores apresentam um bom desempenho. Estes índices são essenciais em estudos de concordância, pois permitem avaliar a consistência das classificações atribuídas por diferentes avaliadores ou por diferentes métodos de medição.

4.5.1 Índice Kappa

O índice Kappa, desenvolvido por Cohen (1960), mede a proporção de concordância entre avaliadores ajustada pela concordância esperada ao acaso. A fórmula para o índice Kappa é dada por:

$$Kp = \frac{P_o - P_e}{1 - P_e} \quad \text{Eq.18}$$

Na Eq. (18), P_o e P_e representam as médias das proporções de concordâncias observadas e esperadas, respectivamente. P_o é a proporção de concordâncias observadas e P_e é a proporção de concordâncias esperadas pelo acaso.

De acordo com os critérios da AIAG (2010), os valores do índice Kappa (K_p) variam entre -1 e 1 , onde um valor de 1 indica uma concordância perfeita, e 0 indica que a concordância é igual à esperada pelo acaso. Um valor de -1 , entretanto, indica uma concordância menor do que a esperada pelo acaso. Este rigor está relacionado ao fato de que, mesmo que um evento seja classificado de forma próxima, qualquer discrepância é considerada um erro. Por exemplo, um evento classificado como extremamente seco com base em estimativas e severamente seco com base em medições de satélite ainda será computado como um erro devido à falta de concordância exata.

4.5.2 Índice Kendall

O coeficiente de Kendall, introduzido por Kendall e Smith (1939), é utilizado para medir o grau de concordância entre diferentes avaliadores. Esse coeficiente também pode ser aplicado para avaliar a consistência interna nas classificações feitas por cada avaliador. O cálculo é realizado por meio das Equações (19–22), onde S representa a soma dos quadrados das somas das classificações R_i , m é o número de bases de dados e k é o número de eventos.

$$Kd = \frac{12S}{m^2k(k^2 - 1)} \quad \text{Eq.19}$$

$$R_i = \sum_{j=1}^m CL_{ij} \quad \text{Eq.20}$$

$$R2 = \frac{1}{k} \sum_{i=1}^k R_i \quad \text{Eq.21}$$

$$S = \sum_{i=1}^k (R_i - R_2)^2 \quad \text{Eq.22}$$

No geral, admite-se que o índice Kappa é mais rigoroso do que o coeficiente de concordância de Kendall. Esse rigor está ligado ao fato de que, ainda que um evento seja extremamente seco (CL1) com base nas estimativas do satélite e severamente seco (CL2) com base nas previsões do modelo, sob o ponto de vista do satélite houve erro ao detectar a classe de severidade do evento. Em outras palavras, já que não houve concordância exata entre as classificações real prevista pelo satélite e a classificação do modelo, computa-se essa situação como erro.

No cálculo do índice de concordância de Kendall, leva-se em consideração que, ainda que os eventos não tenham sido exatamente iguais, suas categorias são próximas. Por exemplo, classificar um evento CL1 como CL2 é considerado menos grave do que classificar um evento extremamente seco como um extremamente úmido. Os valores de Kd variam de 0 a 1, sendo 0 a discordância perfeita entre as bases de dados e 1 a concordância perfeita entre as classificações. Já os valores de Kp podem ser negativos e podem assumir o valor máximo de 1. Quanto mais próximos de 1 forem os valores de Kd e Kp, melhor é o desempenho do modelo em prever os tipos de eventos secos e úmidos ao longo do tempo.

Por fim, após obter a análise de classificação por meio de mapas e o cálculo dos índices, formulam-se gráficos que representam os valores dos índices Kappa e Kendall, a fim de entender o comportamento dos modelos com base na variação espacial da Amazônia Legal.

5 RESULTADOS E DISCUSSÕES

5.1 ANÁLISE DESCRITIVA DOS DADOS DE PRECIPITAÇÃO

As estatísticas descritivas mostram variações significativas na precipitação entre diferentes pixels. Alguns pontos importantes a serem destacados incluem: a média de precipitação varia consideravelmente entre os pixels, de aproximadamente 86,8 mm (pixel 369) a 212,4 mm (pixel 272).

A mediana também apresenta variações significativas, refletindo a natureza não uniforme da precipitação em diferentes áreas. Por exemplo, a mediana de precipitação no pixel 369 é de 58,5 mm, enquanto no pixel 272 é de 158 mm.

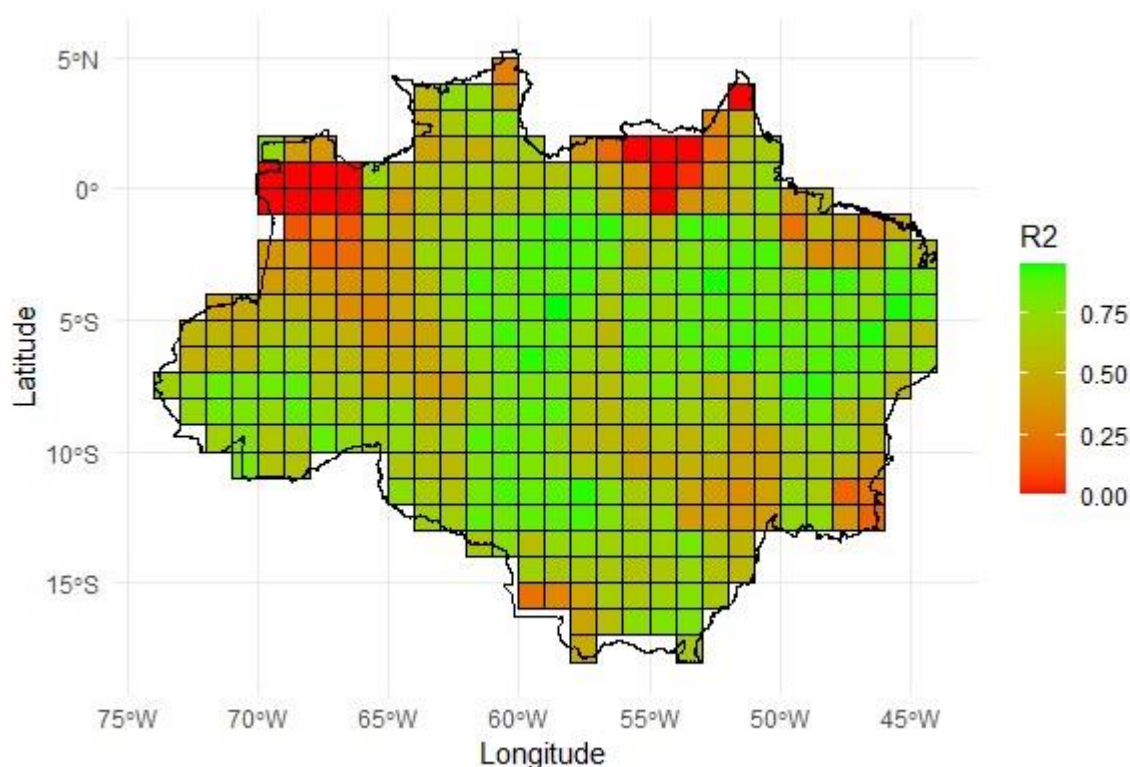
O desvio padrão, que mede a dispersão dos valores em torno da média, varia de 87,1 mm a 164,1 mm. Essa variação indica que alguns pixels experimentam uma variabilidade muito maior na precipitação do que outros. Por exemplo, o pixel 575 tem um desvio padrão de 162,5 mm, indicando grandes variações na precipitação observada.

A precipitação mínima em muitos pixels é 0 mm, indicando períodos sem precipitação. No entanto, a precipitação máxima varia amplamente, de 483 mm (pixel 369) a 961 mm (pixel 272). Esses valores extremos sugerem que alguns pixels estão sujeitos a eventos de precipitação muito intensa.

A análise descritiva revela que há uma variabilidade significativa na precipitação entre os diferentes pixels. As diferenças na média e na mediana da precipitação indicam que algumas áreas recebem consistentemente mais precipitação do que outras. O desvio padrão elevado em alguns pixels indica uma alta variabilidade, o que pode ser um obstáculo na modelagem. Os valores extremos de precipitação mínima e máxima sugerem que certos pixels são mais suscetíveis a eventos de precipitação intensa, o que pode ter implicações importantes para a gestão de recursos hídricos e planejamento de infraestrutura.

5.2 CÁLCULO DO R^2 EM CADA MODELO DESENVOLVIDO

Figura 8 – R^2 dos modelos desenvolvidos pela família ZAGA



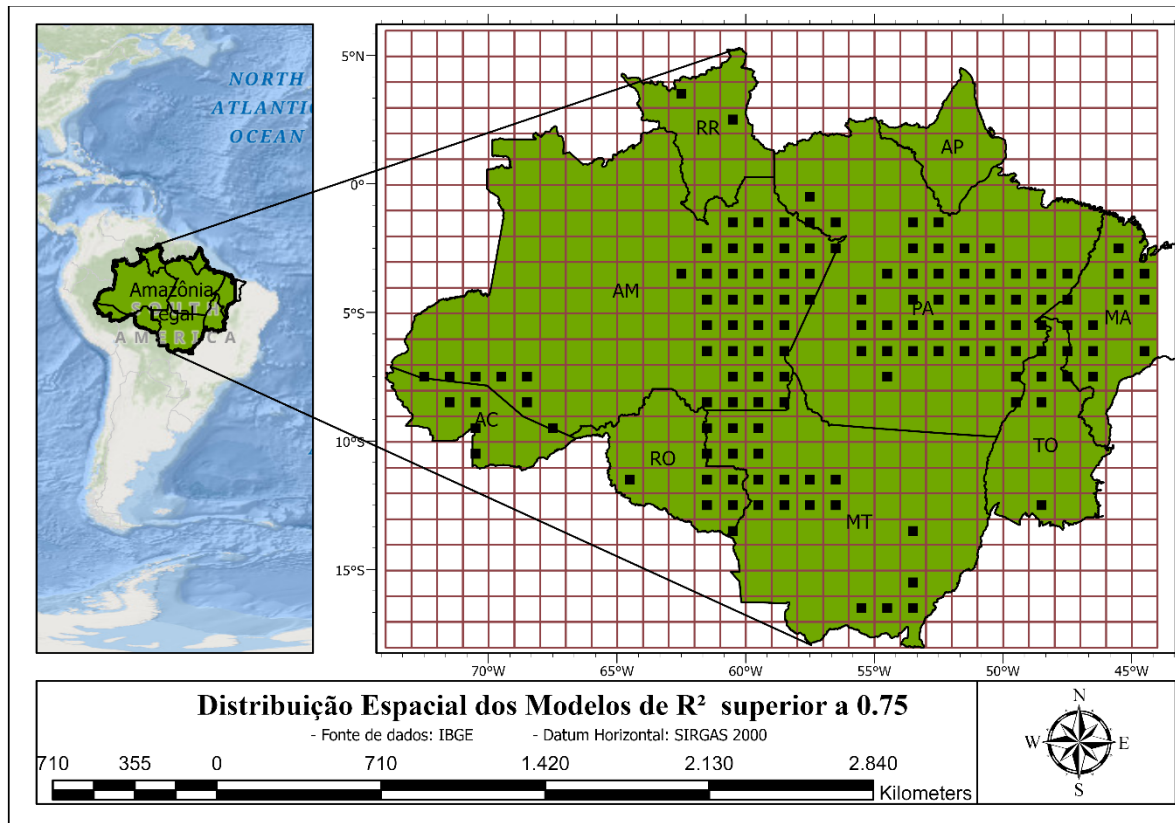
Fonte: Autor

Como resultado da análise do R^2 comparando o acumulo de precipitação entre os valores reais e os preditos dos modelos entre os trimestres: 1) JAN, FEV, MAR; 2) ABR, MAI, JUN; 3) JUL, AGO, SET; 4) OUT, NOV, DEZ. Obtiveram-se os resultados apresentados na Figura 8.

Através da análise espacial dos valores de R^2 dos modelos desenvolvidos utilizando a distribuição ZAGA com R^2 superior a 0.75 (Figura 9), destaca-se que dos 408 pixels modelados, 132 apresentaram valores superiores a 0.75. Dentre estes, observou-se um destaque significativo nos estados do Pará, Maranhão, leste do Amazonas e norte do Tocantins. Na região do Acre e na divisa deste com o Amazonas, também se notou uma aglomeração de modelos que apresentaram desempenho satisfatório. No noroeste do Mato

Grosso e na divisa deste estado com Rondônia, houve igualmente um acúmulo de bons resultados nos modelos construídos. Adicionalmente, alguns pontos de destaque foram identificados no norte do estado de Roraima, no sudeste do Mato Grosso e no sul do Tocantins.

Figura 9 – R^2 de valores superior a 0.75

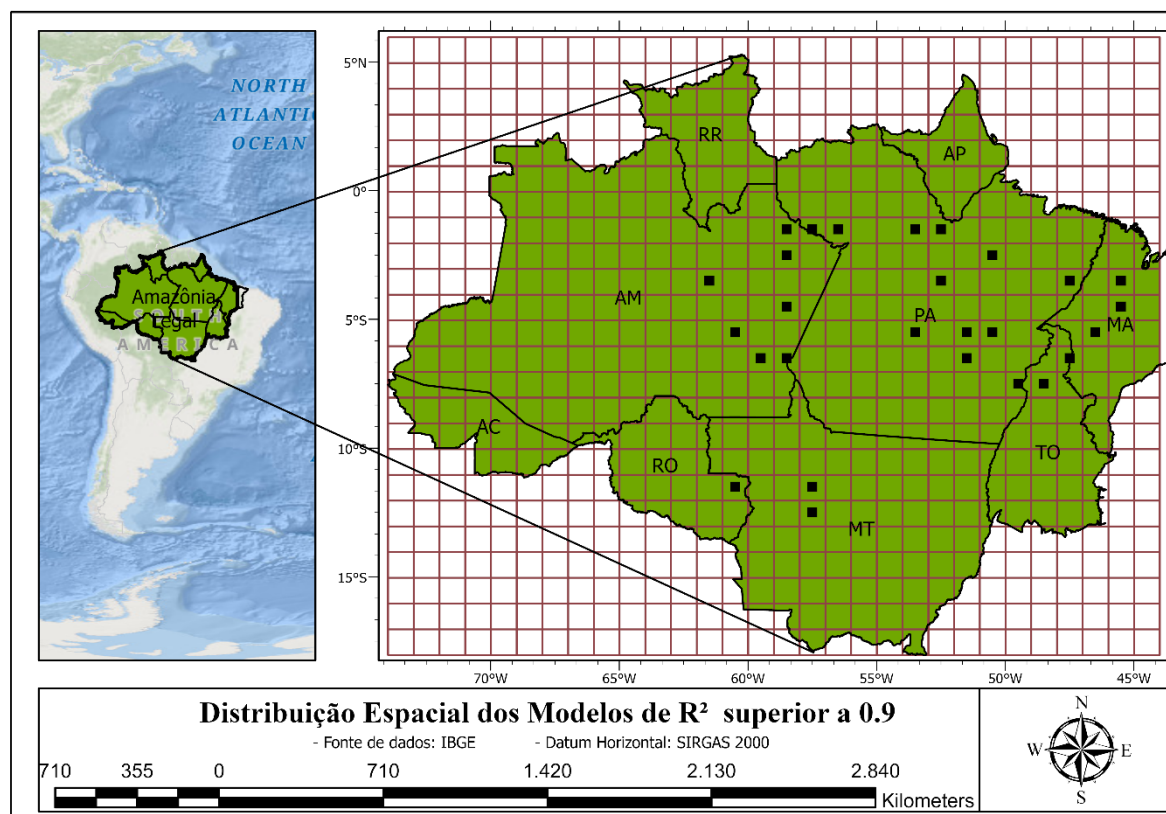


Fonte: Autor

No intuito de encontrar modelos com desempenho mais que satisfatórios, a Figura 10 mapeia apenas os modelos com R^2 superiores a 0.90. Dentre os 408 modelos, 27 apresentaram valores de R^2 acima de 0.9. O estado do Pará se destaca de todos os demais por apresentar 44,44% dos modelos com R^2 superior a 0.9 dentro dos seus limites, seguido pelo Amazonas que apresentou sete grides com modelos em destaque. Apesar de ter sua área reduzida, devido ao escopo do domínio da Amazônia Legal Brasileira, o Maranhão ainda apresentou três grides de destaque, seguido pelo estado do Tocantins, que também apresentou 2 grides destaques, por fim, o estado de Rondônia também apresentou um gride com R^2 superior a 0.9.

As séries temporais de cada pixel representado na Figura 10 está representada na Figura 11, Figura 12, Figura 13. Percebe-se que em todas as distribuições a quantidade de zeros é notável. O modelo inflacionado de zero, também conhecido como a função de distribuição ZAGA, tem se destacado como uma ferramenta robusta para lidar com dados que apresentam uma quantidade significativa de zeros. Desenvolvido por Stasinopoulos *et al.* (2017), este modelo demonstrou ser altamente eficaz na análise de séries temporais com presença de zeros, como indicado por critérios estatísticos comuns, incluindo os Critérios de Informação de Akaike (AIC) e Bayesianos (BIC).

Figura 10 – R^2 de valores superior a 0.90



Fonte: Autor

A presença de períodos frequentes de estiagem nos 27 melhores pixels, muitas vezes resulta em séries temporais que exibem um número substancial de zeros, somados todas as ocorrências de precipitação zero nos 27 melhores modelos foi obtido um total de 407 zeros representados pelos gráficos da Figura 11, Figura 12 e Figura 13. Nesses casos, a distribuição

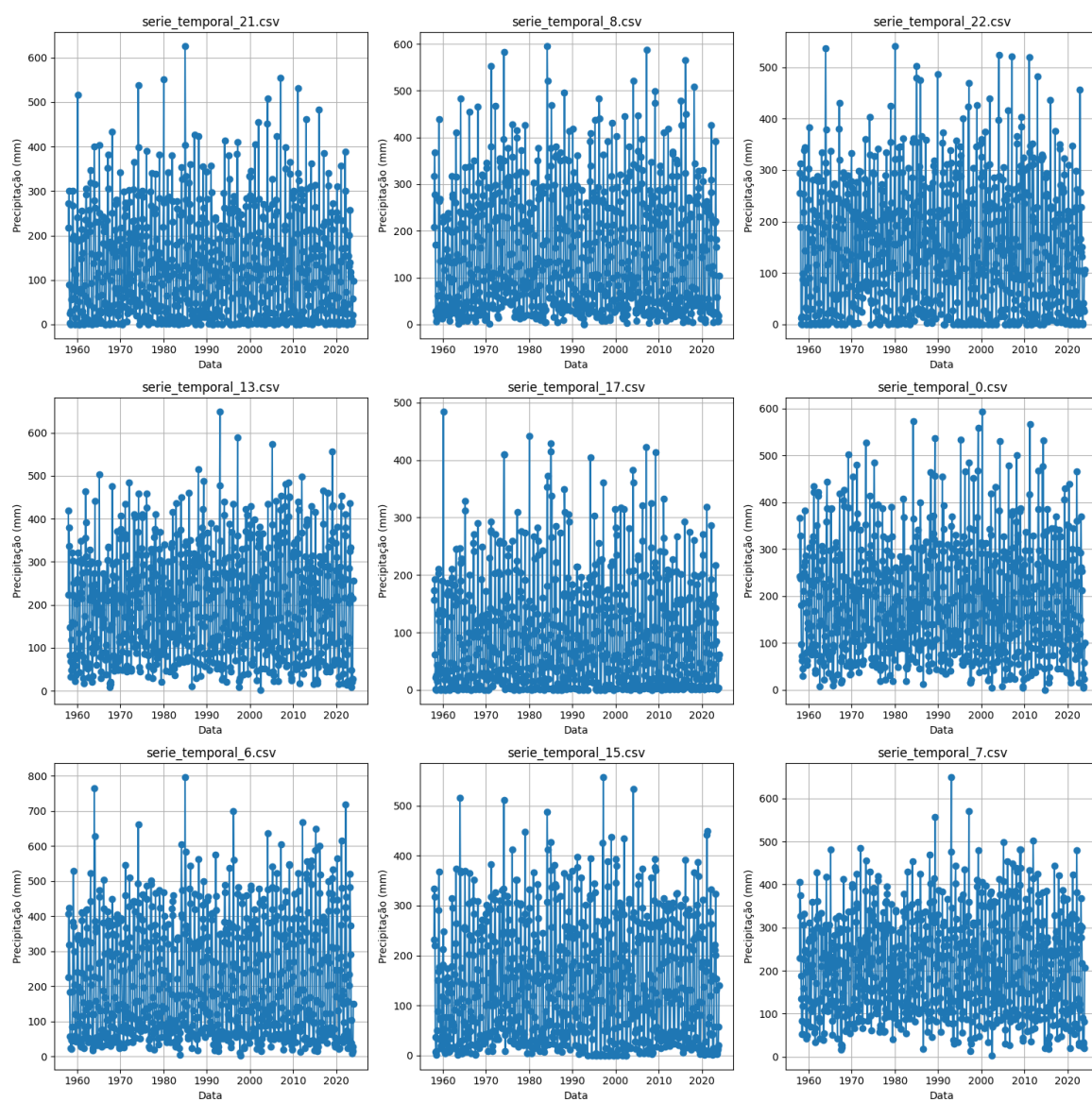
ZAGA se destaca como uma ótima opção para modelar o comportamento dessas séries, proporcionando ajustes precisos e interpretações significativas com modelos de R^2 superiores a 0.9. Em contra partida, em todos os casos em que obtivemos R^2 inferior a 0.10 a quantidade de precipitação nula foi zero. Dessa forma, constata-se que a distribuição ZAGA, não se adequa bem às regiões onde não há estiagem.

Estudos anteriores corroboram a eficácia da distribuição ZAGA na modelagem de séries temporais com presença de zeros. Por exemplo, Rigby e Stasinopoulos (2005), Rashid *et al.* (2016), Rodrigues (2016), Dantas (2020) também observaram resultados positivos ao aplicar essa distribuição em diferentes contextos, reforçando assim sua utilidade e relevância na análise estatística de dados com zeros.

Em cada um dos 27 melhores modelos, a função “stepGAICAll.A” foi capaz de encontrar por meio das análises dos valores dos índices de AIC e BIC a combinação das melhores variáveis para modelar cada um dos pixels. Dessa forma, a Tabela 3 demonstra quantas vezes cada variável explicativa apareceu na modelagem de cada parâmetro estatístico, que no caso da distribuição ZAGA (sigla do inglês *Zero Adjusted Gamma Distribution*) são três parâmetros: locação, escala e assimetria, não se ajustando para curtose.

Ainda nesta tabela, calcula-se a porcentagem de aparição da variável na modelagem de cada parâmetro, é calculado uma média entre eles e por fim, os valores são normalizados com a maior média. Esse valor é apresentado na coluna grau de importância e destaca o AMO, NINO_1_2, PDO, AO e o NINO_4 como as cinco teleconexões que mais influenciaram na modelagem dos modelos.

Figura 11 – Distribuição da precipitação nos 27 melhores pixels – Parte1



Fonte: Autor

Tabela 3 – Importância das anomalias de teleconexões nos 27 melhores modelos

Termos	Parâmetro de Localção (μ)	Parâmetro de Escala Sigma (σ)	Parâmetro de Assimetria (ν)	Grau de Importância
AMO	18,00	5,00	4,00	1,00
NINO_1_2	13,00	9,00	4,00	0,96
PDO	14,00	8,00	1,00	0,85
AO	3,00	10,00	6,00	0,70
NINO4	12,00	2,00	0,00	0,52
NAO	1,00	6,00	3,00	0,37
NINO_3_4	4,00	1,00	1,00	0,22
NINO3	4,00	1,00	1,00	0,22
SOI	0,00	0,00	0,00	0,00

Fonte: Autor

Tabela 4 – Importância da temperatura e vento como variáveis explicativas nos 27 melhores modelos

Termos	Parâmetro de Localção (μ)	Parâmetro de Escala Sigma (σ)	Parâmetro de Assimetria (ν)	Grau de Importância
Temperatura Máxima	27,00	27,00	13,00	1,00
Temperatura Mínima	27,00	24,00	13,00	0,96
Vento	18,00	11,00	4,00	0,49

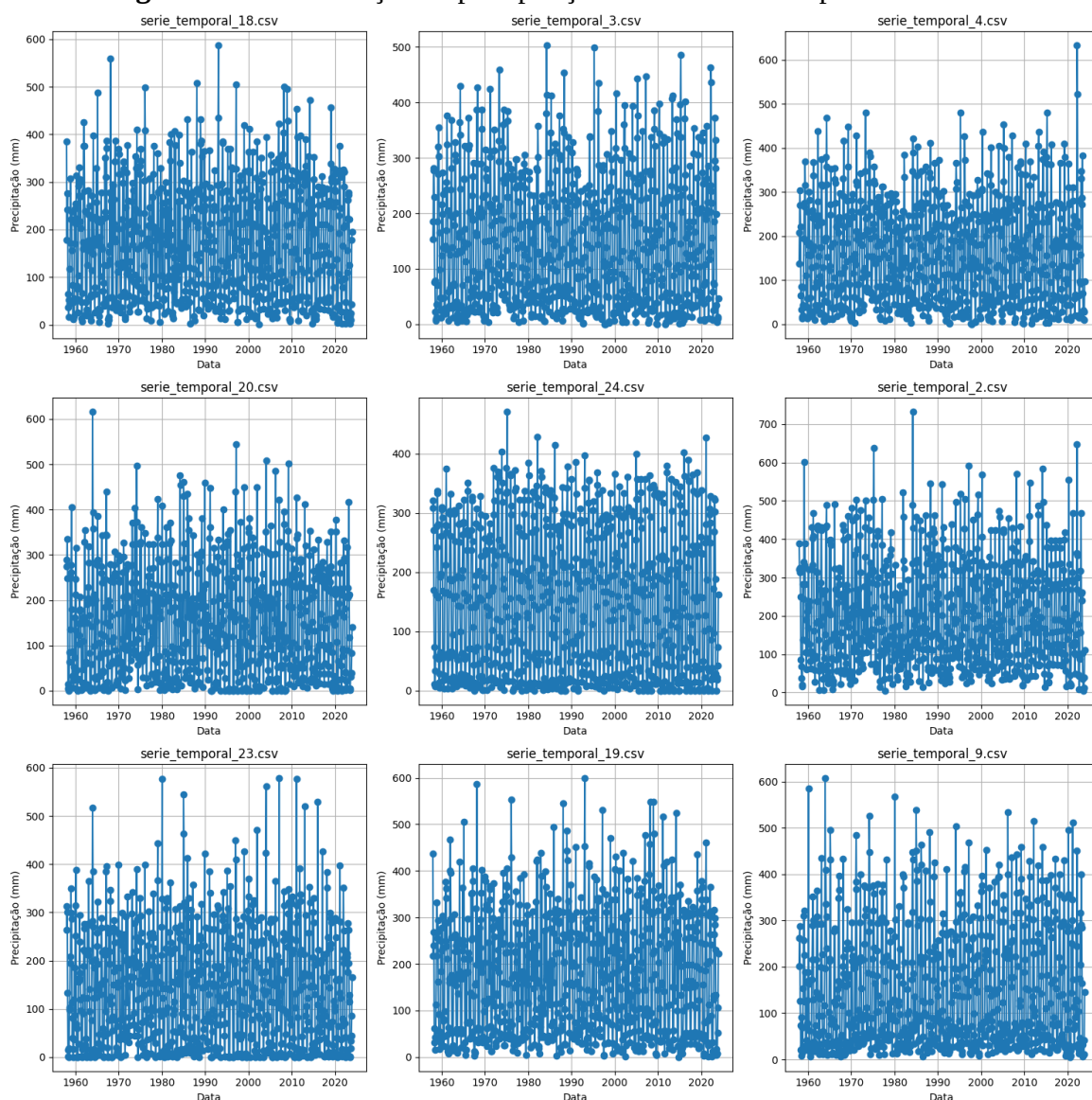
Fonte: Autor

De modo geral, dentre as variáveis que não eram teleconexões, àquelas que mais se destacaram foram a Temperaturas Máxima e a Temperatura Mínima, ambas apresentaram uma forte influência nas modelagens, se enquadrando como uma das variáveis mais importantes das modelagens dos 27 modelos.

Quando analisado a correlação e a tendência das variáveis explicativas podemos notar por meio da Figura 14 que a temperatura máxima demonstrou em todos os 27 pixels uma forte correlação negativa. Já para a temperatura mínima a correlação foi negativa para onze dos vinte e sete pixels, para os outros 16 pixels a correlação se demonstrou positiva com destaque para o pixel 541 que obteve uma alta correlação de 0.72. Da mesma forma que

a temperatura máxima, em todos os vinte e sete modelos analisados a velocidade do vento obteve correlação negativa.

Figura 12 – Distribuição da precipitação nos 27 melhores pixels - Parte 2

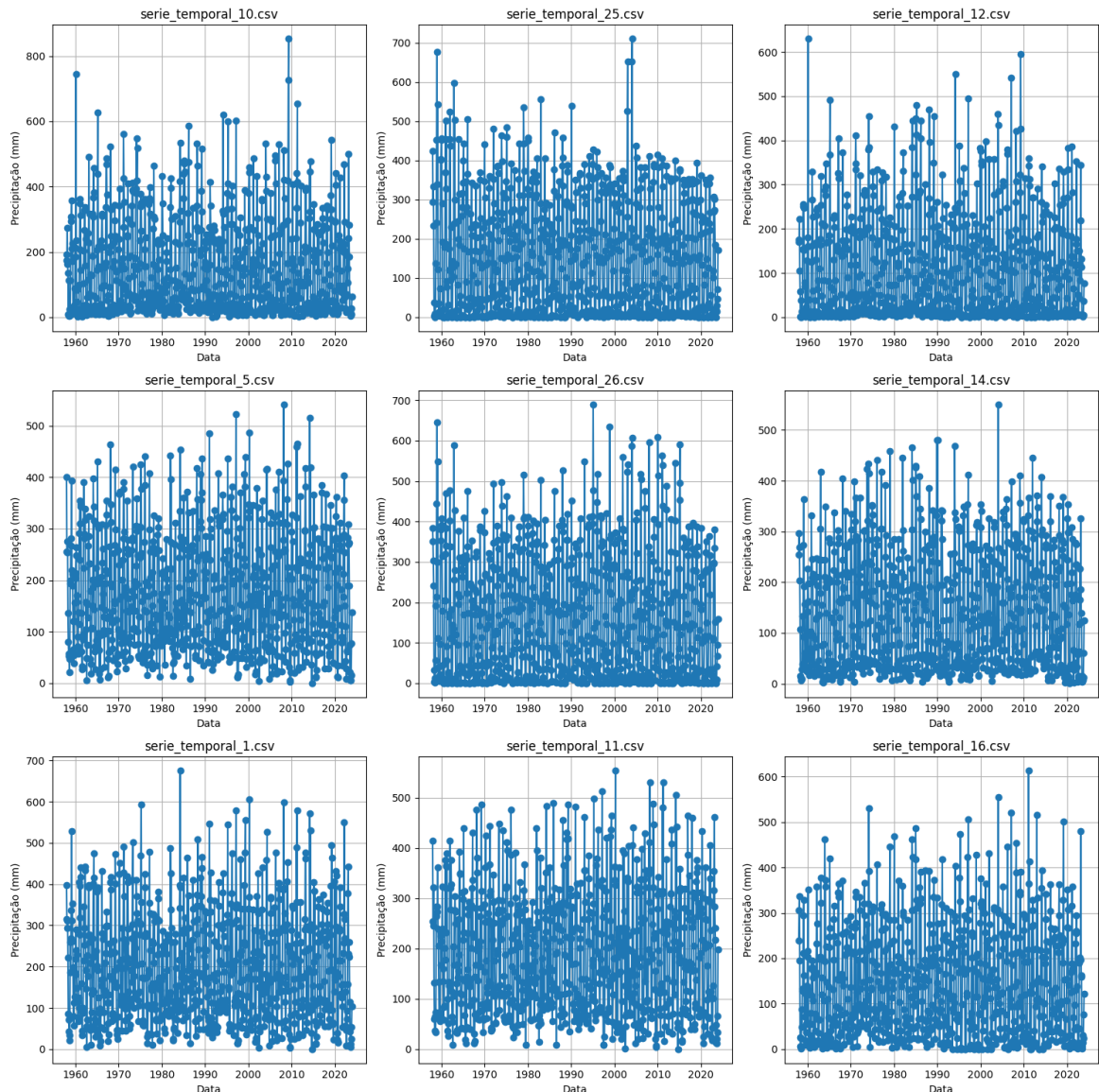


Fonte: Autor

Analisando as variáveis de teleconexões, o AMO se apresentou com uma correlação negativa para todos os grids analisados. Já o NAO demonstrou um comportamento inverso com correlações positivas para todos os modelos. Quinze pixels demonstraram uma relação negativa entre a precipitação e os valores do AO, enquanto o restante obteve valores positivos de correlação. Apesar da forte influência descrita na Tabela 3, o PDO apresentou

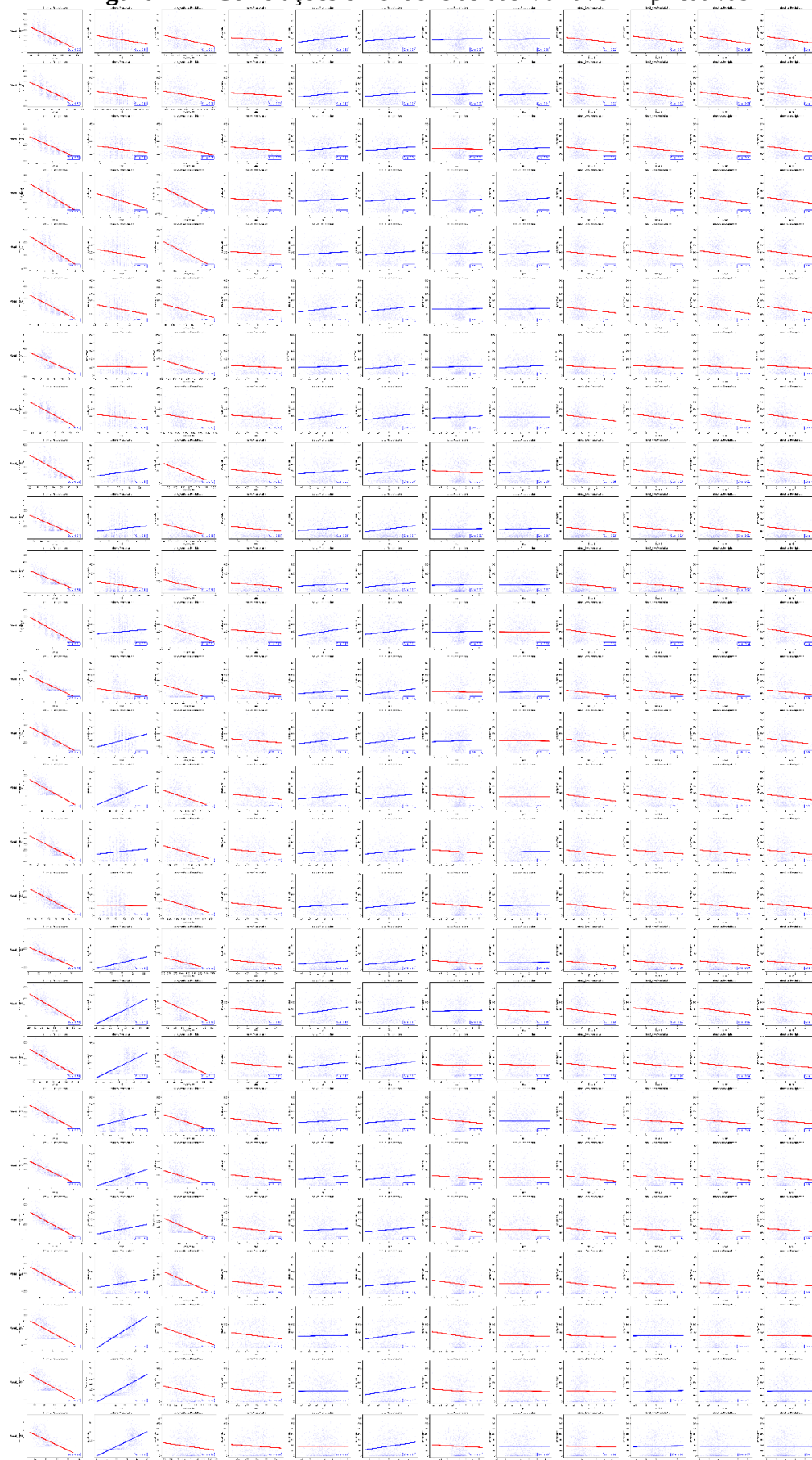
correlações próximas a zero e, em sua grande maioria positivas. Em todos os casos o NINO_1_2 se mostrou inversamente proporcional aos valores de precipitação. Já o NINO_3_4 só obteve três comportamentos de correlação positiva, para os pixels 541,544 e 575. As teleconexões NINO_4 e NINO_3 tiveram um comportamento semelhante, com praticamente todas as correlações sendo negativas, com exceção dos pixels 544 e 575.

Figura 13 – Distribuição da precipitação nos 27 melhores pixels – Parte 3



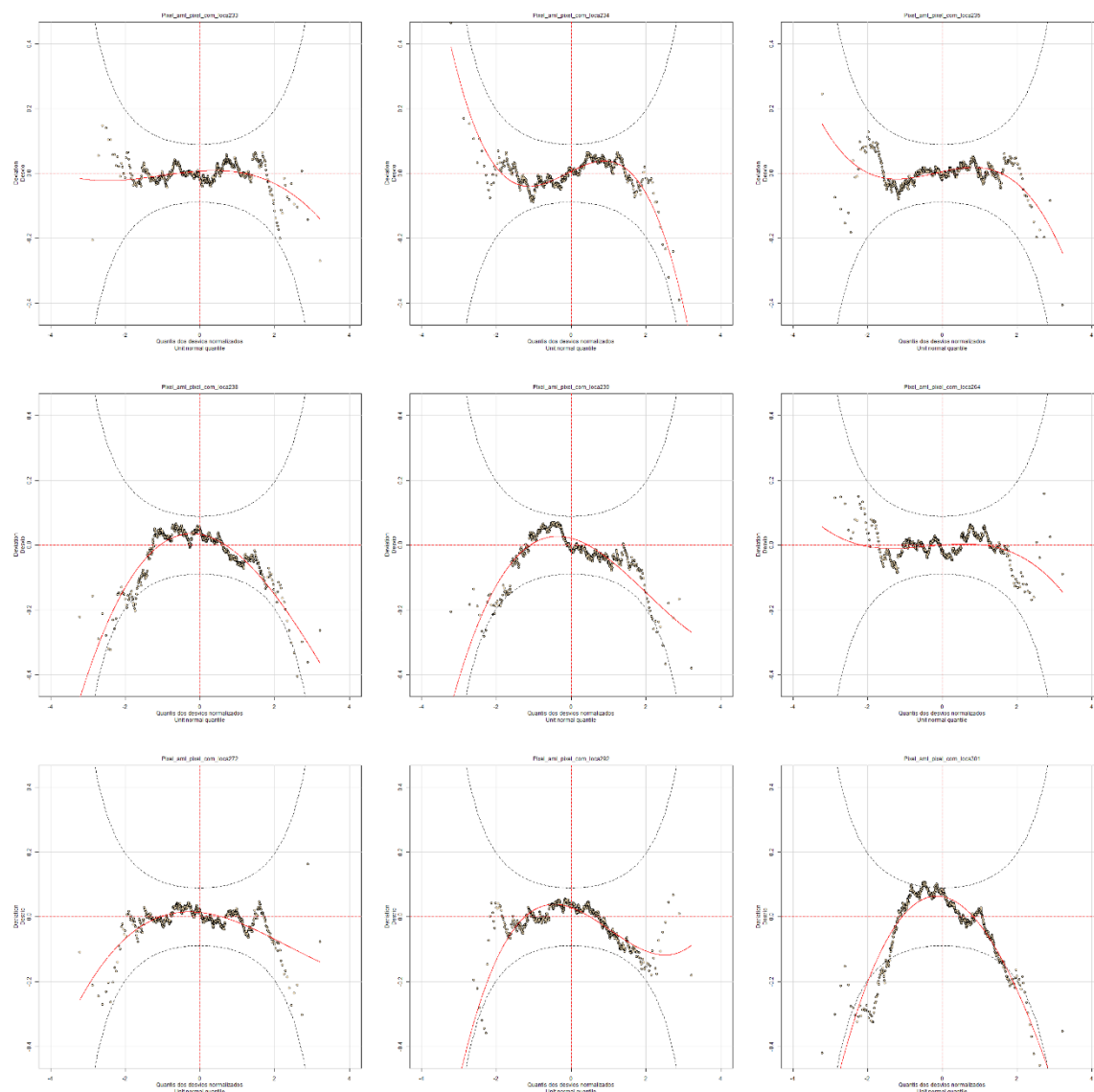
Fonte: Autor

Figura 14 – Correlações e Tendências das Variáveis Explicativas



Analisando as variáveis de teleconexões, o AMO se apresentou com uma correlação negativa para todos os grids analisados. Já o NAO demonstrou um comportamento inverso com correlações positivas para todos os modelos.

Figura 15 – Worm plot dos pixels 233-301



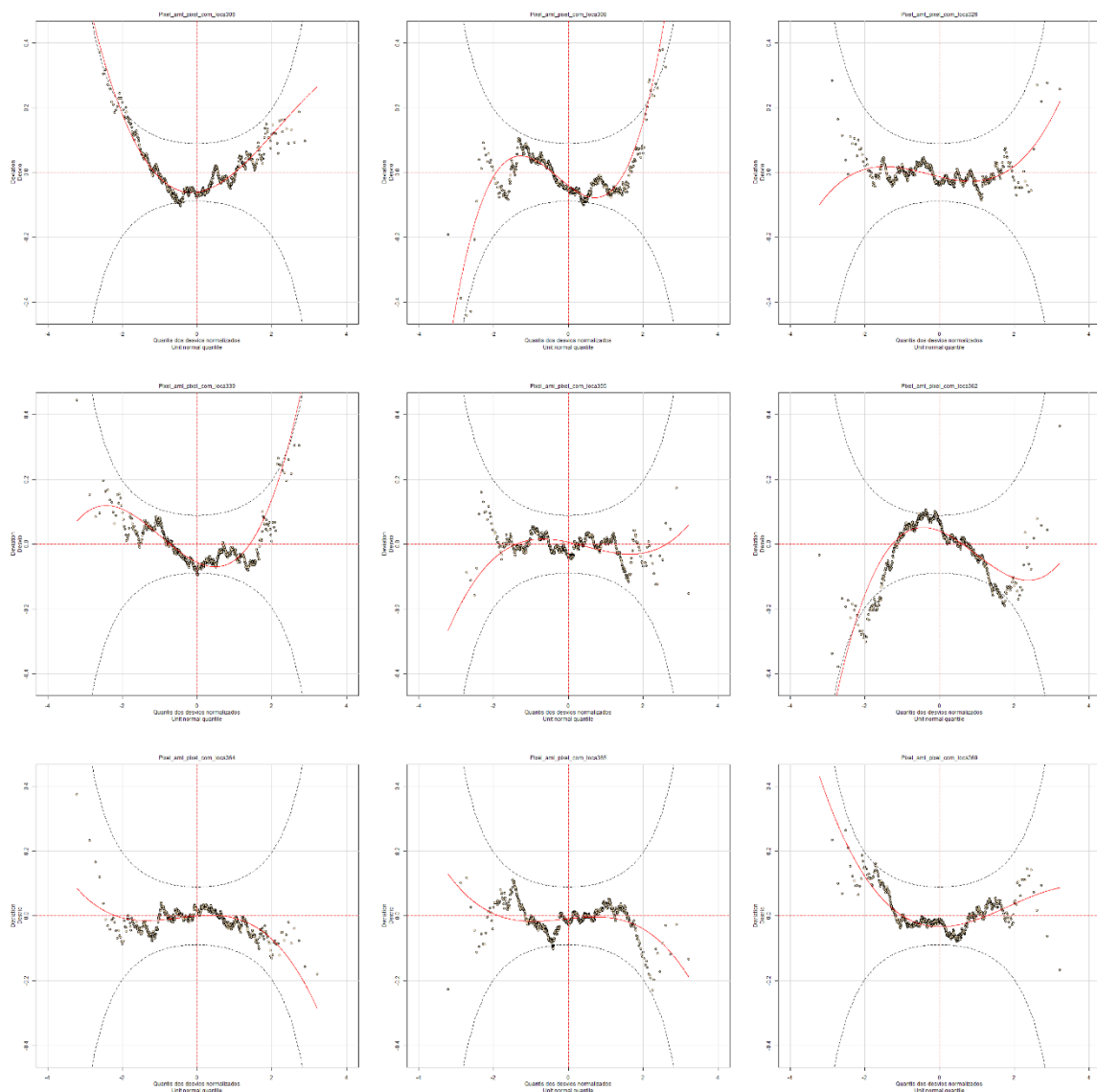
Fonte: Autor

De maneira geral, o comportamento residual quando analisado por meio do worm plot de todos os melhores modelos foram muito bons. Para o Pixel 233, os resíduos estão majoritariamente concentrados em torno da linha zero, com uma leve tendência de desvio nas extremidades. A curtose é relativamente baixa, sugerindo que a distribuição dos resíduos é levemente achatada. A "minhoca" formada pelos pontos segue um padrão quase linear, com pequenas ondulações ao redor da linha central, indicando um bom ajuste nos quantis

centrais. A proximidade dos pontos com a linha horizontal indica que a distribuição dos resíduos está próxima de uma distribuição normal padrão, com poucos pontos fora das elipses de confiança, o que sugere um modelo bem adequado.

No caso do Pixel 234, observa-se uma curtose levemente alta nas extremidades, indicando caudas mais pesadas. Há uma distribuição assimétrica com mais pontos desviando-se à esquerda. A minhoca mostra uma curva acentuada nas extremidades. Conforme a Tabela 1, a inclinação positiva da minhoca sugere que a variância ajustada é muito pequena, e a forma de U invertido indica que a distribuição ajustada é assimétrica à direita.

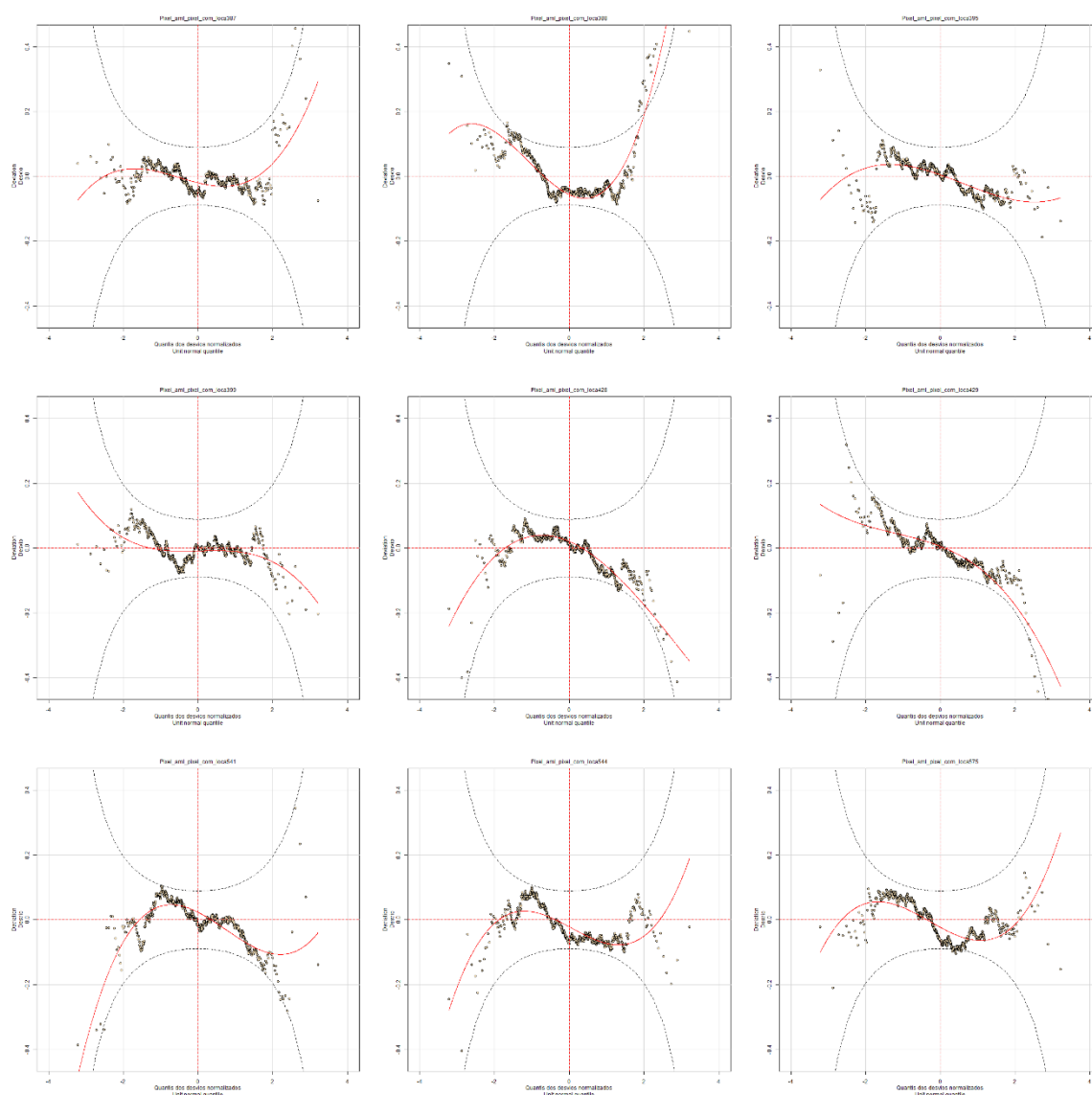
Figura 16 – Worm plot dos pixels 306-369



Fonte: Autor

Para o Pixel 235, os resíduos estão bem centralizados, com uma distribuição levemente assimétrica. A curtose é moderada, com algumas caudas mais pesadas. A minhoca é relativamente linear nos quantis centrais, mas apresenta pequenas curvaturas nas extremidades. Alguns pontos ultrapassam as linhas tracejadas, sugerindo que o modelo pode estar subestimando a variabilidade nos dados extremos. A ideia principal do worm plot reforça que a proximidade dos pontos à linha horizontal indica uma distribuição adequada dos resíduos nos quantis centrais.

Figura 17 – Worm plot dos pixels 387-575



Fonte: Autor

No caso do Pixel 238, a distribuição dos resíduos é levemente achatada, com uma curtose baixa. Os pontos estão bem concentrados em torno da linha central. A minhoca apresenta uma curvatura simétrica, permanecendo dentro das linhas tracejadas na maioria dos casos, indicando um bom ajuste geral. Pequenos desvios nas extremidades sugerem a necessidade de ajustes menores no modelo. A tabela indica que a forma de U da minhoca aponta para uma distribuição ajustada assimétrica à esquerda. Segundo a interpretação do worm plot, a maioria dos pontos está dentro dos intervalos de confiança de 95%, indicando um modelo adequado com poucas necessidades de ajustes.

Para o Pixel 239, a curtose é baixa nos quantis centrais, com caudas mais pesadas nas extremidades. A distribuição é levemente assimétrica. A minhoca mostra uma curvatura significativa nas extremidades, com alguns pontos fora das linhas tracejadas. Isso sugere que o modelo está capturando bem os quantis centrais, mas pode precisar de ajustes para os extremos. A tabela sugere que a forma em U invertido da minhoca indica uma distribuição assimétrica à direita. Conforme a ideia principal do worm plot, a presença de pontos fora das elipses de confiança sugere que o modelo pode ser melhorado para capturar a variabilidade dos resíduos nas extremidades.

No caso do Pixel 264, a curtose é moderada, com resíduos bem distribuídos em torno da linha zero. A distribuição é relativamente simétrica. A minhoca é quase linear nos quantis centrais, mas apresenta curvatura nas extremidades. Poucos pontos ultrapassam as linhas tracejadas, indicando um bom ajuste geral com pequenos desvios. A interpretação do worm plot reforça que, com a maioria dos pontos dentro dos intervalos de confiança, o modelo está bem ajustado.

Para o Pixel 272, a distribuição dos resíduos é levemente assimétrica, com uma curtose moderada. Os pontos estão concentrados principalmente em torno da linha central. A minhoca apresenta uma leve curvatura, com a maioria dos pontos dentro das linhas tracejadas. Pequenos desvios nas extremidades sugerem uma boa adequação do modelo, com necessidade de ajustes menores. A forma levemente curva da minhoca pode sugerir uma leve assimetria na distribuição ajustada.

Para o Pixel 292, a curtose é moderada, com uma distribuição levemente assimétrica. Os resíduos estão concentrados em torno da linha central. A minhoca segue uma trajetória levemente curva, com alguns pontos nas extremidades ultrapassando as linhas tracejadas.

Isso indica que o modelo está adequado para os quantis centrais, mas pode precisar de ajustes para os dados extremos. A tabela sugere que a forma de U invertido da minhoca indica uma distribuição assimétrica à direita. A interpretação do worm plot reforça que a proximidade dos pontos à linha horizontal indica um ajuste adequado, mas com necessidade de melhorias nas extremidades.

No caso do Pixel 301, a curtose é alta, com caudas pesadas e uma distribuição assimétrica. Os resíduos mostram uma variação considerável nas extremidades. A minhoca apresenta uma curvatura acentuada, com vários pontos fora das linhas tracejadas. Isso sugere que o modelo não está capturando de maneira perfeita a variabilidade dos dados extremos. A tabela sugere que a forma em S da minhoca, curvada para baixo, indica que as caudas da distribuição ajustada são muito leves. De toda forma, são poucos pontos que ultrapassam as hipérboles de confiança.

Para o Pixel 306, a curtose é elevada, indicando caudas pesadas. A distribuição dos resíduos é assimétrica, com mais desvios à esquerda. A minhoca mostra uma curva significativa, especialmente nas extremidades, com alguns pontos ultrapassando as linhas tracejadas. Isso indica que o modelo não está capturando bem a variabilidade nos quantis extremos. A tabela sugere que a forma em S da minhoca, curvada para cima, indica que as caudas da distribuição ajustada são muito pesadas. Entretanto, o ajuste ainda se mostra com um bom resultado.

Para o Pixel 308, a análise mostra que os resíduos estão concentrados ao redor da linha zero, mas com desvios significativos nas extremidades, indicando uma alta curtose. A "minhoca". A ideia principal do worm plot reforça que, a maioria dos pontos estão dentro das elipses de confiança, apenas um número não tão considerável de pontos está fora delas, indicando a presença de outliers, mas de um bom ajuste.

No caso do Pixel 326, os resíduos estão bem centralizados em torno da linha zero, com uma distribuição relativamente simétrica e uma curtose moderada. A minhoca segue uma trajetória quase linear, mas com uma leve curva. A proximidade dos pontos com a linha horizontal indica que a distribuição dos resíduos está próxima de uma distribuição normal padrão, e a maioria dos pontos está dentro das elipses de confiança, sugerindo um modelo muito bem adequado.

Para o Pixel 339, a análise revela uma distribuição assimétrica dos resíduos, com caudas mais pesadas, indicando uma alta curtose. A minhoca apresenta uma forma em S invertido, indicando que as caudas da distribuição ajustada são muito leves. A presença de poucos pontos fora das elipses de confiança sugere que o modelo está capturando relativamente bem a variabilidade dos resíduos, com um pouco de dificuldade nas extremidades.

No caso do Pixel 355, os resíduos estão bem centralizados, com uma distribuição levemente assimétrica e uma curtose moderada. A minhoca segue uma trajetória quase linear. Todos os pontos estão dentro das elipses de confiança, indicando um ajuste geral adequado.

Para o Pixel 362, a análise mostra que os resíduos têm uma distribuição assimétrica, com caudas pesadas, indicando uma alta curtose. A minhoca apresenta uma forma em U invertido, sugerindo que a distribuição ajustada é assimétrica à direita. A presença de alguns pontos fora das elipses de confiança indica que o modelo pode estar subestimando a variabilidade dos resíduos, especialmente nas extremidades, entretanto, de maneira geral o ajuste do modelo é bom.

No caso do Pixel 364, os resíduos estão bem centralizados em torno da linha zero, com uma distribuição relativamente simétrica e uma curtose moderada. A minhoca segue uma trajetória quase linear. Todos os pontos estão dentro das elipses de confiança, sugerindo um ajuste geral perfeito.

Para o Pixel 365, a análise revela uma distribuição assimétrica dos resíduos, com caudas mais pesadas, indicando um grau de curtose maior que o caso anterior. A minhoca apresenta uma forma em S invertido, indicando que as caudas da distribuição ajustada são muito leves. A presença de poucos pontos fora das elipses de confiança sugere que o modelo está capturando bem a variabilidade dos resíduos.

No caso do Pixel 369, a análise mostra que os resíduos têm uma distribuição assimétrica, com caudas pesadas, indicando uma curtose moderada. A minhoca apresenta uma forma em U invertido, sugerindo que a distribuição ajustada é assimétrica à direita. Apenas um ponto ficou fora das hipérboles de confiança, caracterizando um bom comportamento para o modelo analisado.

Para o Pixel 387, os resíduos estão bem centralizados em torno da linha zero, com uma distribuição relativamente simétrica e uma curtose moderada. A minhoca segue uma trajetória quase linear, mas com uma leve curva, indicando que a variância ajustada pode ser ligeiramente pequena. A maioria dos pontos está dentro das elipses de confiança, sugerindo um ajuste geral adequado, com pequenas inconsistências nas extremidades.

No caso do Pixel 388, a análise revela uma distribuição assimétrica dos resíduos, com caudas mais pesadas, indicando uma alta curtose. A minhoca apresenta uma forma em S invertido, indicando que as caudas da distribuição ajustada são muito leves. A presença de alguns pontos fora das elipses de confiança sugere que o modelo pode não estar capturando tão bem a variabilidade dos resíduos, contudo isso acontece mais nas extremidades e, de maneira geral, os pontos estão dentro das hipérboles de confiança.

Para o Pixel 395, os resíduos estão bem centralizados em torno da linha zero, com uma distribuição relativamente simétrica e uma curtose moderada. A "minhoca" formada pelos pontos segue uma trajetória quase linear. De acordo com a Tabela 1, a inclinação positiva da minhoca sugere que a variância ajustada é muito pequena. Todos os pontos estão dentro das elipses de confiança, indicando um ajuste geral perfeito.

No caso do Pixel 399, os resíduos têm uma distribuição levemente assimétrica, com caudas mais pesadas, indicando uma curtose moderada. A minhoca apresenta uma forma quase linear. Conforme a Tabela 1, a inclinação da minhoca sugere que a variância ajustada não é muito grande. Não há presença de pontos fora das elipses de confiança, o que indica um bom resultado para o pixel.

Para o Pixel 428, a análise revela uma distribuição assimétrica dos resíduos, com caudas mais pesadas, indicando uma alta curtose. A minhoca apresenta uma forma de U invertido, indicando que a distribuição ajustada é assimétrica à direita. Segundo a Tabela 1, a inclinação negativa sugere que a variância ajustada é um pouco acentuada. Contudo, há poucos pontos fora das elipses de confiança o que sugere um bom ajuste.

No caso do Pixel 429, os resíduos estão bem centralizados em torno da linha zero, mas com desvios significativos nas extremidades, indicando uma alta curtose. A minhoca apresenta uma forma de U invertido com uma inclinação negativa, sugerindo que há certa variância ajustada, além disso é notável que a distribuição ajustada é assimétrica à direita. A

ideia principal do worm plot reforça que, a maioria dos pontos estão dentro das elipses de confiança, indicando a presença de poucos outliers e de um bom ajuste.

Para o Pixel 541, a análise mostra que os resíduos têm uma distribuição assimétrica, com caudas pesadas, indicando uma alta curtose. A minhoca apresenta uma forma de U invertido, sugerindo que a distribuição ajustada é assimétrica à direita. Há apenas três pontos fora das hipérboles de confiança o que demonstra um bom ajuste do modelo.

No caso do Pixel 544, os resíduos estão bem centralizados em torno da linha zero, com uma distribuição relativamente simétrica e uma curtose moderada. A minhoca segue uma trajetória quase linear, mas com uma leve inclinação positiva, indicando que a variância ajustada pode ser ligeiramente pequena. Todos os pontos estão dentro das elipses de confiança, sugerindo um ótimo ajuste geral.

Para o Pixel 575, a análise revela uma distribuição um pouco assimétrica dos resíduos, com caudas mais pesadas, indicando uma certa curtose. A minhoca apresenta uma forma de U invertido, indicando que a distribuição ajustada é assimétrica à direita. A presença de poucos pontos fora das elipses de confiança sugere que o modelo está capturando bem a variabilidade dos resíduos.

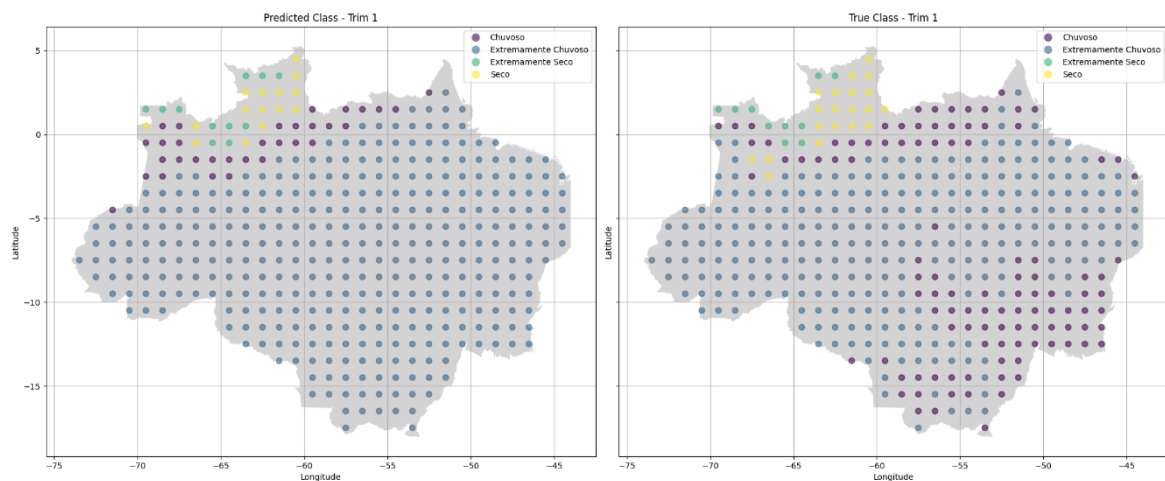
5.3 ANÁLISE QUALITATIVA DO MODELO

Cada pixel tem seu valor predito por meio da soma da precipitação em cada trimestre. Dessa forma, através das Figura 18, Figura 19, Figura 20, Figura 21 podemos avaliar o comportamento da classificação da chuva através da comparação da classificação do valor predito por cada modelo de cada pixel e sua classificação real correspondente.

No primeiro trimestre (Figura 18), observa-se uma correspondência razoável entre as classes previstas e as classes verdadeiras. As regiões do norte e noroeste apresentam uma maior correspondência com a classificação "Extremamente Seco" e "Seco". No entanto, há algumas discrepâncias nas áreas centrais e sul, onde as classes previstas indicam "Extremamente Chuvoso", mas a classe real indica "Chuvoso", entretanto, o resultado para o primeiro trimestre se demonstrou bastante satisfatório.

No segundo trimestre (Figura 19), as classes previstas e verdadeiras mostram uma melhor correspondência geral em comparação com o primeiro trimestre. A região norte e noroeste continua a mostrar uma boa correspondência com as classificações "Extremamente Chuvoso" e "Chuvoso". As regiões centrais e sul também mostram uma correspondência melhor, mas ainda há algumas discrepâncias, especialmente em áreas próximas a Latitude - 5 ° em que há alguns erros entre a classificação “Chuvoso” e “Seco”.

Figura 18 – Resultado da análise qualitativa para o 1º Trimestre



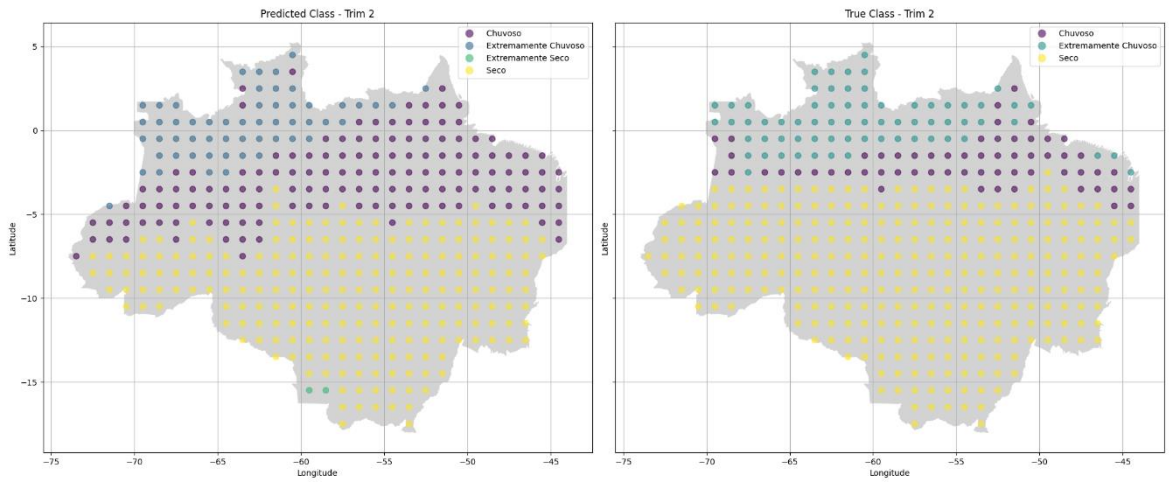
Fonte: Autor

O terceiro trimestre (Figura 20) a grande maioria dos pontos são “Extremamente Seco” o que facilita para um bom comportamento do modelo que se utiliza da ZAGA como família de distribuição. Na parte noroeste houveram alguns pontos que não corresponderam com algumas classificações “Chuvosa”, quando na verdade seria a classificação “Seca”.

No quarto trimestre (Figura 21), a correspondência entre as classes previstas e verdadeiras não foi tão boa quanto aos outros trimestres. A região ao sul da Latitude -5° apresentou bom comportamento, porém, com alguns erros de classificação na parte sudeste onde classificou alguns pontos como “Chuvoso”, mas deveria ter sido “Extremamente Chuvoso”. Perto desta latitude também houveram alguns erros majoritariamente na

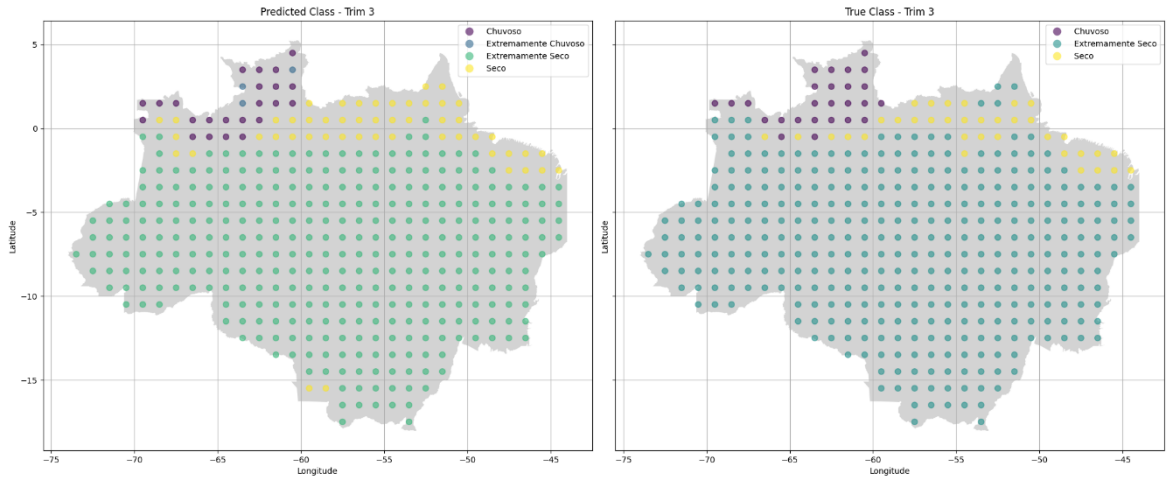
classificação “Chuvoso” como “Seco”. De maneira geral, o modelo classificou de forma precisa o comportamento da chuva para o quarto trimestre.

Figura 19 – Resultado da análise qualitativa para o 2º Trimestre



Fonte: Autor

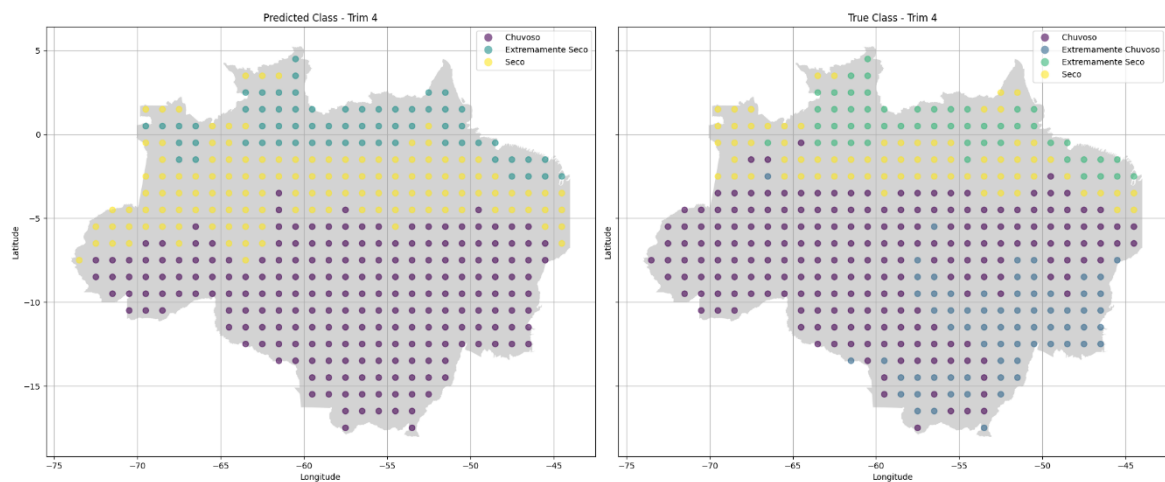
Figura 20 – Resultado da análise qualitativa para o 3º Trimestre



Fonte: Autor

O desempenho do modelo de previsão de classes de precipitação na região amazônica variou ao longo dos quatro trimestres, mostrando uma ótima correspondência entre as classes previstas e as reais. O modelo teve um desempenho satisfatório no primeiro trimestre, melhorou no segundo trimestre, apresentou um bom comportamento no terceiro trimestre devido à predominância de uma única classe ("Extremamente Seco"), mas teve mais dificuldades no quarto trimestre, com alguns erros de classificação. De maneira geral, o modelo demonstrou eficácia na previsão das classes de precipitação

Figura 21 – Resultado da análise qualitativa para o 4º Trimestre



Fonte: Autor

O índice Kappa apresentado na Figura 22 e Figura 23 é uma medida estatística que corrige a concordância observada para o acaso, variando de 0 a 1, onde 0 indica nenhuma concordância além do acaso e 1 indica concordância perfeita. A interpretação dos valores de Kappa no contexto da avaliação de modelos de previsão é crucial para entender a eficácia e limitações do modelo.

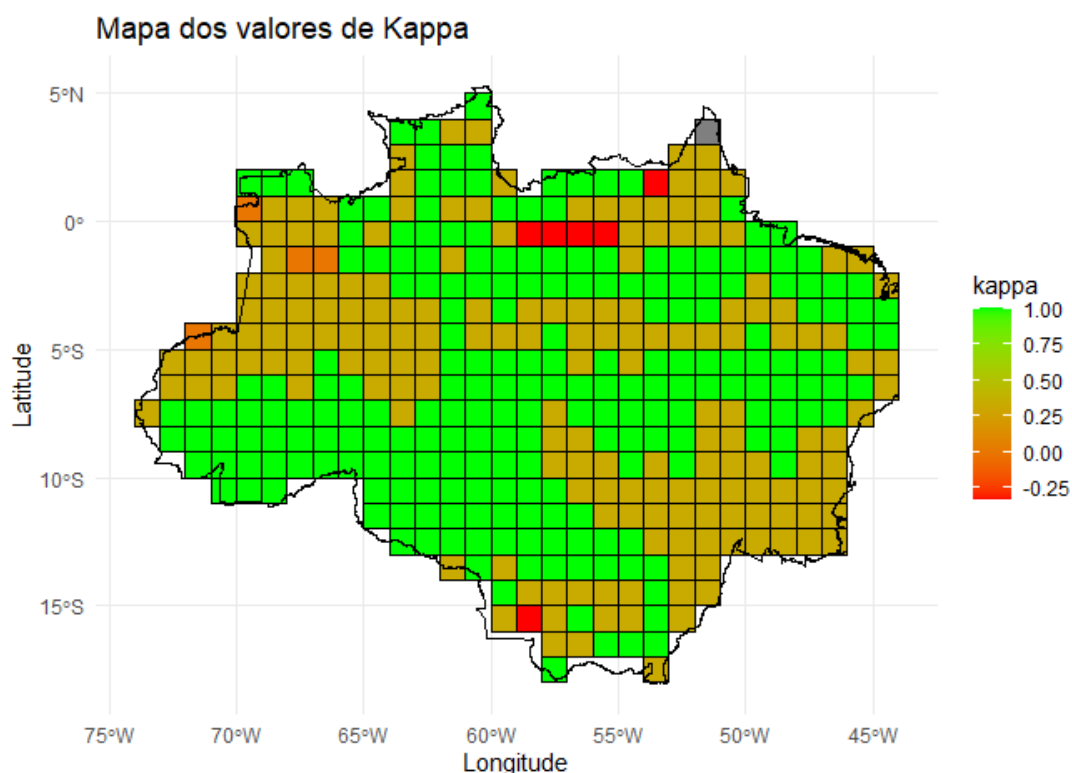
Os resultados do mapa revelam uma variação espacial significativa na performance do modelo de previsão de precipitação. As regiões com valores altos de Kappa, indicados pela cor verde, mostram uma forte concordância entre as classificações previstas e as observadas. Notavelmente, estas áreas incluem o centro-norte e noroeste da Amazônia Legal, onde o índice Kappa próximo a 1 sugere que o modelo é altamente eficaz na previsão das

classes de precipitação. Este desempenho robusto pode ser atribuído a características climáticas mais homogêneas e previsíveis nessas regiões.

Por outro lado, áreas com valores moderados de Kappa, representadas por tons de amarelo e laranja (valores entre 0.50 e 0.75), são observadas predominantemente no noroeste e no sudeste da Amazônia Legal. A concordância moderada nesta região indica que, embora o modelo consiga capturar algumas tendências de precipitação, ainda há variações que ele não prevê adequadamente.

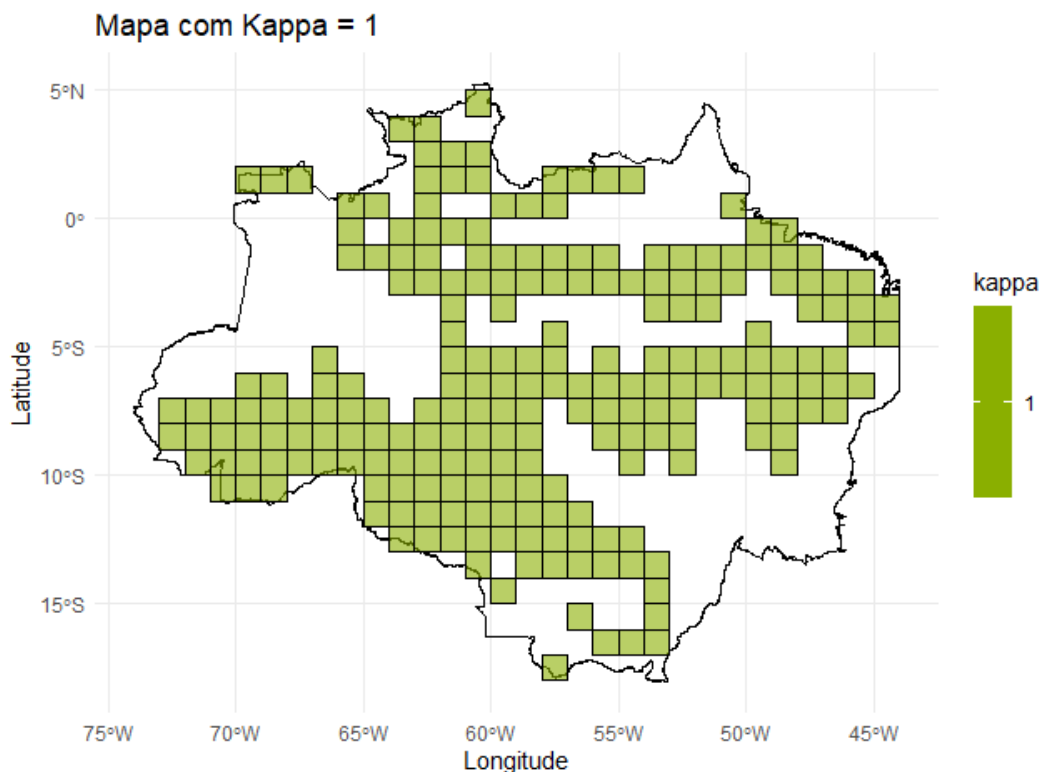
As regiões com valores baixos de Kappa, destacadas em vermelho, principalmente no norte da Amazônia Legal, apresentam uma concordância muito baixa entre as previsões do modelo e os valores reais. Nestes locais, a baixa eficácia do modelo pode ser atribuída a diversos fatores, como a heterogeneidade climática local, insuficiência de dados históricos ou inadequações no modelo que não capturam a variabilidade específica dessas áreas, bem como a modelagem com a utilização de uma família de distribuição não tão adequada para a região em específico.

Figura 22 – Valores de Kappa



Fonte: Autor

Figura 23 – Valores de Kappa igual a 1



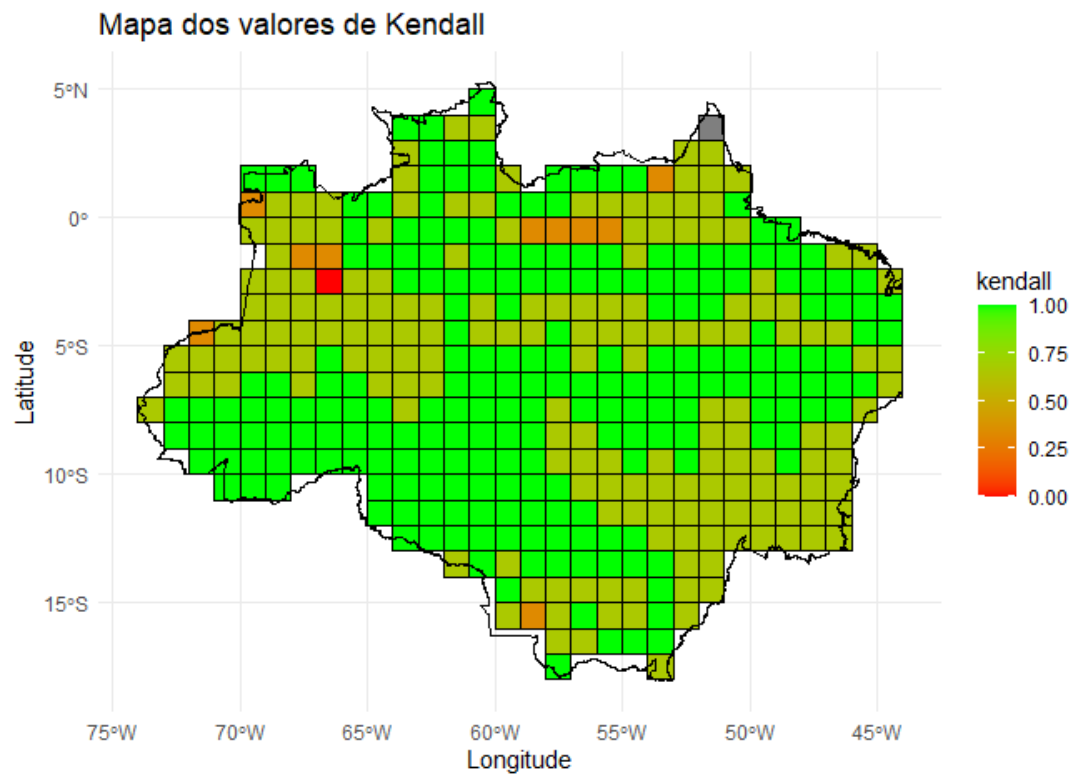
Fonte: Autor

A análise dos mapas gerados a partir do índice de concordância de Kendall (Figura 24, Figura 25), aplicado para comparar os valores de classificação previstos e reais, fornece informações valiosas sobre o comportamento do modelo em diferentes regiões da Amazônia Legal brasileira. O índice de Kendall é uma medida estatística que avalia a associação entre duas classificações, variando de -1 (discordância total) a 1 (concordância total).

A Figura 24 ilustra uma distribuição contínua dos valores de Kendall, permitindo uma visão mais detalhada da performance do modelo em toda a Amazônia Legal Brasileira. As áreas em verde representam alta concordância, enquanto as áreas em tons de vermelho e laranja indicam menor concordância. As regiões de alta concordância estão nas partes centrais e algumas áreas ao norte e sudoeste da Amazônia Legal. As regiões de Moderada Concordância indicam que, embora o modelo tenha um bom desempenho, há alguma variação entre os valores previstos e reais. Essas áreas incluem partes do Pará e Mato Grosso.

Poucas áreas tiveram baixa concordância, podendo ser encontrado alguns pixels ao norte e ao leste da região, mostram baixa concordância.

Figura 24 – Valores de Kendall

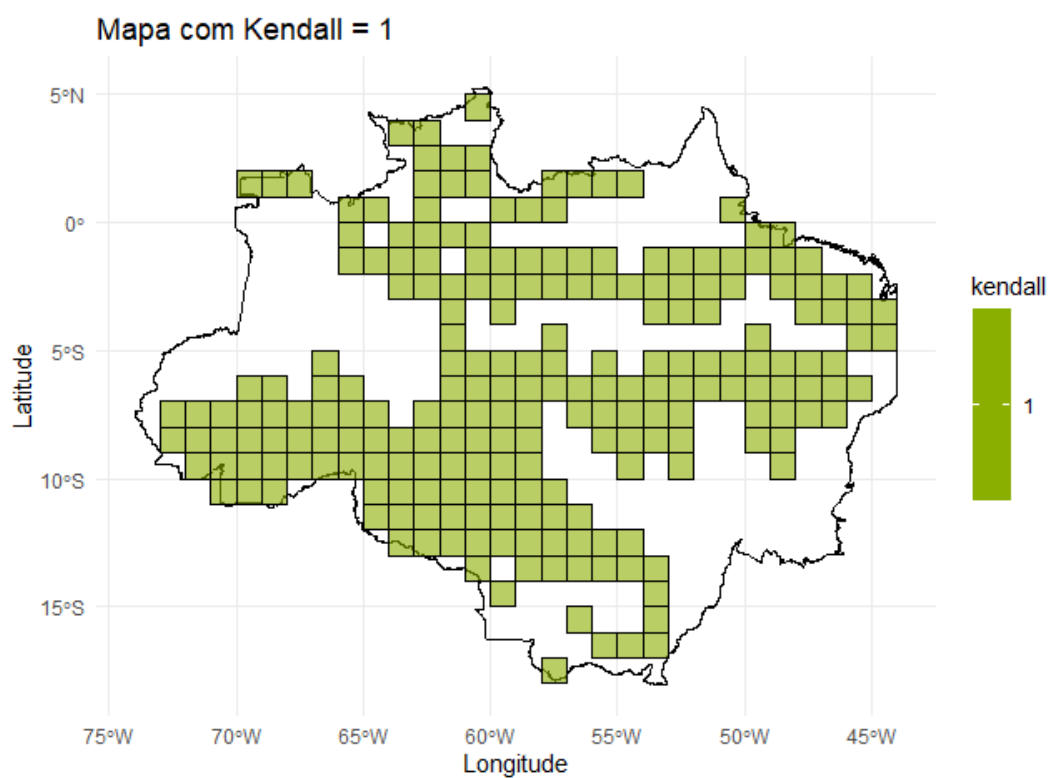


Fonte: Autor

Na Figura 25, tem-se as regiões onde o índice de Kendall é igual a 1, observamos áreas onde há concordância perfeita entre os valores previstos e reais. As regiões destacadas em verde representam locais onde o modelo conseguiu prever com precisão as classes de precipitação para todos os trimestres analisados. Essas áreas estão dispersas por diversas partes da Amazônia Legal, incluindo o Norte da Amazônia Legal Brasileira, o Centro, o Sudoeste e Sudeste da Amazônia: Partes do Acre, Rondônia e Mato Grosso, bem como algumas áreas no Pará, apresentam concordância perfeita. A presença dessas regiões com alto índice de concordância indica que o modelo é bastante eficaz em prever as classes de precipitação nessas áreas.

Os mapas de Kendall, juntamente com a análise visual, indicam que o modelo tem um desempenho muito bom na previsão das classes de precipitação para a maior parte da Amazônia Legal.

Figura 25 – Valores de Kendall iguais a 1



Fonte: Autor

6 CONCLUSÕES

Este estudo teve como objetivo modelar funções que explicam o comportamento da precipitação na Amazônia Legal, considerando o comportamento das anomalias das teleconexões, temperaturas máximas e mínimas e a velocidade do vento. Para avaliar a eficácia dos modelos de previsão de precipitação, foram utilizados os índices de concordância de Kendall e Kappa. A análise foi realizada por trimestres do ano de 2021, e as previsões dos modelos foram agrupadas e analisadas com base nesses índices. Além disso, análises estatísticas e residuais foram feitas nos modelos mais eficazes, destacando-se os melhores grides com os melhores modelos.

Os resultados, apresentados na Figura 6, demonstram que os modelos baseados na distribuição ZAGA alcançaram valores de R^2 superiores a 0,75 em 132 dos 408 pixels modelados. Este desempenho foi particularmente notável nos estados do Pará, Maranhão, leste do Amazonas e norte do Tocantins, além de áreas específicas no Acre, Rondônia e Mato Grosso. A Figura 7 corrobora esses achados, indicando que uma parte significativa dos modelos desenvolvidos apresentou desempenho satisfatório. Além disso, conforme ilustrado na Figura 8, 27 pixels alcançaram um R^2 superior a 0,90, com o estado do Pará abrigando 44,44% desses modelos de alto desempenho, seguido pelo Amazonas, Maranhão e Rondônia. Esses resultados evidenciam a robustez da distribuição ZAGA em certas regiões da Amazônia Legal.

A análise temporal dos melhores pixels, ilustrada nas Figuras 9, 10 e 11, reforçou a eficácia do modelo inflacionado de zero (ZAGA) na captura de séries temporais com alta incidência de zeros, característica comum em dados de precipitação. Por outro lado, em casos onde obtivemos R^2 inferior a 0,10, a quantidade de precipitação nula foi zero, indicando que a distribuição ZAGA não se adapta bem a regiões sem estiagem.

A função “stepGAICAll.A” foi utilizada para identificar a combinação ideal de variáveis explicativas em cada um dos 27 melhores modelos, com base nos índices de AIC e BIC. Analisou-se a frequência com que cada variável explicativa apareceu na modelagem desses parâmetros e calculou-se a média dessas frequências, que foram então normalizadas para determinar o grau de importância de cada variável.

Os resultados destacaram as teleconexões AMO, NINO_1_2, PDO, AO e NINO_4 como as cinco variáveis mais influentes na modelagem dos 27 melhores modelos. O AMO

apresentou a maior importância, seguido por NINO_1_2 e PDO. Variáveis como NINO_3_4 e SOI mostraram menor influência. A análise das correlações das variáveis de teleconexão mostrou que o AMO apresentou uma correlação negativa consistente em todos os grids analisados, enquanto o NAO teve correlações positivas.

A correlação negativa consistente do AMO com a precipitação na Amazônia Legal sugere que, durante fases positivas do AMO (temperaturas mais altas), há uma tendência de redução da precipitação na região. Esse fenômeno pode ser atribuído a mudanças na circulação atmosférica que afetam os padrões de precipitação.

Embora o NAO esteja mais diretamente associado ao clima da Europa e da América do Norte, suas flutuações podem ter impactos indiretos na Amazônia Legal. A correlação positiva observada entre a NAO e a precipitação na Amazônia pode ser explicada pela influência dessa teleconexão na circulação atmosférica global, que pode alterar a posição e a intensidade das células de alta e baixa pressão em outras partes do mundo, incluindo a América do Sul. Durante fases positivas do NAO, pode ocorrer uma intensificação dos ventos de oeste para leste, afetando a circulação tropical e, potencialmente, aumentando a convergência de umidade na região amazônica.

A variabilidade das teleconexões influenciou de maneira diversa as previsões de precipitação. Por exemplo, NINO_1_2 mostrou-se inversamente proporcional aos valores de precipitação em todos os casos, enquanto NINO_3_4 apresentou apenas três correlações positivas, destacando os pixels 541, 544 e 575. NINO_4 e NINO_3 exibiram comportamentos semelhantes, com correlações predominantemente negativas, exceto nos pixels 544 e 575.

Entre as variáveis que não eram teleconexões, as temperaturas máxima e mínima se destacaram significativamente. A análise de correlação revelou uma correlação negativa da temperatura máxima em todos os 27 pixels analisados, enquanto a temperatura mínima apresentou correlação negativa em onze pixels e positiva nos outros dezesseis. A velocidade do vento também mostrou correlação negativa em todos os modelos.

A correlação positiva entre a temperatura mínima e a precipitação, observada na maioria dos casos, pode ser explicada por vários fatores. Nas situações de correlações negativas, o arrefecimento noturno pode reduzir a convecção e a formação de nuvens pela manhã, impactando negativamente a precipitação. Além disso, em regiões de alta elevação,

temperaturas mínimas mais baixas podem estar associadas a condições mais secas devido à topografia e às variações de temperatura local. Por outro lado, as correlações positivas podem resultar da maior retenção de calor e umidade durante a noite, o que favorece a formação de nuvens e precipitação no dia seguinte. Em áreas com densa cobertura vegetal, temperaturas mínimas mais altas podem indicar a presença de nuvens noturnas que retêm calor, sugerindo umidade disponível para precipitação.

A correlação negativa entre a temperatura máxima e a precipitação pode ser atribuída a diversos fatores. Primeiro, a evapotranspiração: em condições de alta temperatura máxima, a taxa de evapotranspiração aumenta, o que retira umidade do solo e da vegetação, reduzindo a disponibilidade de umidade para a formação de nuvens e precipitação. Segundo, a convecção térmica: embora altas temperaturas máximas possam inicialmente favorecer a convecção e a formação de nuvens de tempestade, se persistirem por um período prolongado, isso pode indicar a presença de alta pressão predominante, que está associada a condições mais secas. Terceiro, secas e ondas de calor: a correlação negativa também pode refletir períodos de seca ou ondas de calor, onde altas temperaturas coincidem com baixa precipitação.

A análise qualitativa dos resultados do modelo por trimestre mostrou que, em janeiro, fevereiro e março, houve uma correspondência razoável entre as classes previstas e verdadeiras, especialmente nas regiões norte e noroeste, com as classificações "Extremamente Seco" e "Seco". No entanto, ocorreram discrepâncias nas áreas centrais e sul, onde o modelo previu "Extremamente Chuvoso" em vez de "Chuvoso". No segundo trimestre, a correspondência melhorou nas regiões norte e noroeste, mas ainda houve erros nas áreas próximas à Latitude -5° , oscilando entre "Chuvoso" e "Seco".

No terceiro trimestre, a maioria dos pontos foi corretamente classificada como "Extremamente Seco", embora alguns pontos no noroeste não tenham sido correspondidos corretamente. No quarto trimestre, a correspondência foi menos precisa, com erros significativos na parte sudeste, onde "Chuvoso" foi previsto em vez de "Extremamente Chuvoso". Apesar desses erros, o modelo classificou muito bem o comportamento da chuva.

A análise dos índices Kappa e Kendall forneceu uma visão detalhada da performance espacial dos modelos. Os valores de Kappa variaram significativamente, indicando áreas de alta concordância localizadas por toda a Amazônia legal, exceto no noroeste e sudeste.

A análise dos resíduos através de worm plots forneceu percepções adicionais sobre a qualidade do ajuste dos modelos. Conforme ilustrado nas Figuras 15, 16 e 17, os worm plots dos pixels analisados indicaram que a maioria dos modelos apresentou um excelente ajuste nos quantis centrais, embora algumas variações tenham sido notadas nas extremidades.

Os resultados deste estudo reforçam a eficácia dos modelos baseados na distribuição ZAGA para a previsão de precipitação em regiões com alta variabilidade climática, como a Amazônia Legal. A robustez do modelo é particularmente evidente em regiões com alta incidência de zeros nas séries temporais. No entanto, em regiões com menor variabilidade e sem estiagem, a performance do modelo pode ser inferior, o que indica a necessidade de adaptações ou a utilização de modelos com outras famílias de distribuições.

Em suma, este estudo contribui significativamente para a compreensão da performance dos modelos de previsão de precipitação na Amazônia Legal, destacando a utilidade do GAMLSS e da distribuição ZAGA. Ele também fornece uma base sólida para futuras pesquisas na área.

REFERÊNCIAS BIBLIOGRÁFICAS

- AKAIKE, H. Information theory and an extension of the maximum likelihood principle. In: Selected Papers of Hirotugu Akaike. [S.l.]: Springer, 1998. p. 199–213.
- ALCÂNTARA, Cássio de. **Modelos aditivos generalizados para posição, escala e forma (GAMLSS) na modelagem da área miocárdica sob risco de necrose**. 2018. Trabalho de Conclusão de Curso (Bacharelado em Estatística) – Universidade Federal de Uberlândia, Uberlândia, 2018.
- ALMEIDA, F. A.; OLIVEIRA-JÚNIOR, J. F.; CUBO, P. Spatiotemporal Rainfall and Temperature Trends throughout the Brazilian Legal Amazon, 1973-2013. *International Journal of Climatology*, v. 36, n. 5, p. 2233-2247, 2016. DOI: 10.1002/joc.4505.
- ANDREOLI, R. V.; KAYANO, M. T.; SOUZA, R. A. F.; SILVA, J. C. S. Seasonal precipitation prediction for Brazil using canonical correlation analysis. *Climate Dynamics*, v. 49, n. 11, p. 4135-4149, 2017. DOI: 10.1007/s00382-017-3553-7.
- Anyah, R.O.; Forootan, E.; Awange, J.; Khaki, M. Understanding linkages between global climate indices and terrestrial water storage changes over Africa using GRACE products. *Sci. Total Environ.* 2018, 635, 1405–1416.
- Bana, R.S. 2014. Agrotechniques for conserving water and sustaining production in rainfed agriculture. *Indian Farming* 63(10): 30–35.
- Bana, R.S., Rana, K.S., Dass, A., Choudhary, A.K., Pooniya, V., Vyas, A.K., Kaur, R., Sepat, S. and Rana, D.S. 2013. A manual on dryland farming and watershed management, Division of Agronomy, IARI, New Delhi. pp. 104.
- BAWA, Kamaljit S.; SEIDLER, Reinmar. Deforestation and sustainable mixed-use landscapes: a view from the eastern Himalaya1. **Annals of the Missouri Botanical Garden**, v. 100, n. 3, p. 141-149, 2015.
- Bayer, D.M.; Castro, N.M.R. Modelagem e Previsão de Vazões Médias Mensais do Rio Potiribu Utilizando Modelos de Séries Temporais. *RBRH* 2012, 17, 229–239
- Bezak, N., Šraj, M., Mikoš, M., 2016. Copula-based IDF curves and empirical rainfall thresholds for flash floods and rainfall-induced landslides. *J. Hydrol.* 541, 272–284. <https://doi.org/10.1016/j.jhydrol.2016.02.058>.
- BIRSAN, M. V. *et al.* Streamflow trends in Switzerland. *Journal of Hydrology*, v. 314, n. 1–4, p. 312–329, 2005.
- BIRSAN, M. V. *et al.* Streamflow trends in Switzerland. *Journal of Hydrology*, v. 314, n. 1–4, p. 312–329, 2005.
- BRANDÃO, A. M. Clima Urbano e Enchentes na Cidade do Rio de Janeiro. In: GUERRA, A. J.T.; CUNHA, S. B. da. (orgs.) *Impactos Ambientais Urbanos no Brasil*. 3ª Ed. Rio de Janeiro, Bertrand Brasil, 2005.
- BRIDGMAN, H. A.; OLIVER, J. E. The global climate system: patterns, processes, and teleconnections. Cambridge University Press, 2014.

BRIDGMAN, Howard A.; OLIVER, John E. **The global climate system: patterns, processes, and teleconnections**. Cambridge University Press, 2014.

Bui, D.T., Tsangaratos, P., Ngo, P.-T.T., Pham, T.D., Pham, B.T., 2019. Flash flood susceptibility modeling using an optimized fuzzy rule based feature selection technique and tree based ensemble methods. *Sci. Total Environ.* 668, 1038–1054. <https://doi.org/10.1016/j.scitotenv.2019.02.422>.

Bui, K.-T.T., Tien Bui, D., Zou, J., Van Doan, C., Revhaug, I., 2018. A novel hybrid artificial intelligent approach based on neural fuzzy inference model and particle swarm optimization for horizontal displacement modeling of hydropower dam. *Neural Comput. e Applic.* 29, 1495–1506. <https://doi.org/10.1007/s00521-016-2666-0>.

BUUREN, S. v.; FREDRIKS, M. Worm plot: a simple diagnostic device for modelling growth reference curves. *Statistics in medicine*, Wiley Online Library, v. 20, n. 8, p. 1259–1277, 2001.

BUUREN, S. van. Worm plot to diagnose fit in quantile regression. *Statistical Modelling*, v.7, n.4, p.363–376, 2007.

CANCHALA, Teresita *et al.* Monthly rainfall anomalies forecasting for southwestern Colombia using artificial neural networks approaches. **Water**, v. 12, n. 9, p. 2628, 2020.

CASTRO, A.L.C.; CALHEIROS, L. B. Manual de medicina de desastres. Ministério da Integração Nacional. Secretaria Nacional de Defesa Civil. Brasília: MI, 2007.

CAYAN, D. R.; REDMOND, K. T.; RIDDLE, L. G. ENSO and Hydrologic Extremes in the Western United States. *Journal of Climate*, v. 12, n. 9, p. 2881-2893, 1998. DOI: 10.1175/1520-0442(1999)012<2881>

Chen, J.L.; Wilson, C.R.; Tapley, B.D. The 2009 exceptional Amazon flood and interannual terrestrial water storage change observed by GRACE. *Water Resour. Res.* 2010, 46, W12526.

COLLISCHONN, W.; DORNELLES, F. (2015). Hidrologia para engenharias e ciências ambientais. Associação Brasileira de Recursos Hídricos – ABRH. 2ª Edição. Porto Alegre, 350p.

COLLISCHONN, W.; DORNELLES, F. Hidrologia para engenharias e ciências ambientais. Associação Brasileira de Recursos Hídricos – ABRH, 2ª Edição, Porto Alegre, 2015.

DANTAS, Leydson G. *et al.* Rainfall Prediction in the State of Paraíba, Northeastern Brazil Using Generalized Additive Models. **Water**, v. 12, n. 9, p. 2478, 2020.

DEMÉTRIO, C.; CORDEIRO, G. Modelos lineares generalizados. Simpósio de Estatística Aplicada à Experimentação Agronômica, v.12, 2007.

DE BASTIANI, F.; RIGBY, R. A.; CYSNEIROS, F. J. A.; URIBE-OPAZO, M. A. On empirical quantile functions for modelling the skewness and kurtosis in the GAMLSS framework. *Journal of Applied Statistics*, v. 45, n. 15, p. 2721-2739, 2018.

Deo, R.C., Şahin, M., 2015. Application of the artificial neural network model for prediction of monthly standardized precipitation and evapotranspiration index using hydrometeorological parameters and climate indices in eastern Australia. *Atmos. Res.* 161–162, 65–81. <https://doi.org/10.1016/j.atmosres.2015.03.018>.

Deo, R.C., Salcedo-Sanz, S., Carro-Calvo, L., Saavedra-Moreno, B., 2018. Chapter 10 - Drought prediction with standardized precipitation and evapotranspiration index and support vector

regression models. In: Samui, P., Kim, D., Ghosh, C. (Eds.), Integrating Disaster Science and Management. Elsevier, pp. 151–174. <https://doi.org/10.1016/B978-0-12-812056-9.00010-5>.

DOS SANTOS, D. C. Índices de extremos climáticos baseados na precipitação pluvial e influência de teleconexões no Nordeste do Brasil. Dissertação (Mestrado em Engenharia Civil e Ambiental) - Universidade Federal da Paraíba, João Pessoa, 2018.

DUNN, Peter K.; SMYTH, Gordon K. Randomized quantile residuals. **Journal of Computational and graphical statistics**, v. 5, n. 3, p. 236-244, 1996.

EVANS, A. D.; BENNETT, J. M.; EWENZ, C. M. South Australian rainfall variability and climate extremes. *Climate Dynamics*, v. 33, n. 4, p. 477–493, 2009.

EVANS, A. D.; BENNETT, J. M.; EWENZ, C. M. South Australian rainfall variability and climate extremes. *Climate Dynamics*, v. 33, n. 4, p. 477–493, 2009.

Faccini, F., Luino, F., Paliaga, G., Sacchini, A., Turconi, L., de Jong, C., 2018. Role of rainfall intensity and urban sprawl in the 2014 flash flood in Genoa City, Bisagno catchment (Liguria, Italy). *Appl. Geogr.* 98, 224–241. <https://doi.org/10.1016/j.apgeog.2018.07.022>.

Farhangi M., Kholghi M. e Chavoshian S.A. 2016. Rainfall Trend Analysis of Hydrological Subbasins in Western Iran. *Journal of Irrigation and Drainage Engineering*, 142(10), 1–11. doi: 10.1061/(ASCE)IR.1943-4774.0001040

FLORENCIO, L. Engenharia de avaliações com base em modelos GAMLSS. 2010. 125 p. 2010. Dissertação (Mestrado em Estatística). Departamento de Estatística, Universidade Federal de Pernambuco.

FOLLAND, C. K.; SCOTT, P. A.; KNIGHT, J. R.; GRIMES, D. I. F.; ROWELL, D. P.; THORPE, R. The roles of external forcing and natural oscillations in recent developments of the European winter climate. *Climate Dynamics*, v. 32, n. 1, p. 255-272, 2009. DOI: 10.1007/s00382-008-0445-6.

Gao, L.; Huang, J.; Chen, X.; Chen, Y.; Liu, M. Risk of extreme precipitation under non-stationary conditions during the second flood season in the southeastern coastal region of China. *J. Hydrometeorol.* 2017, 18, 669– 681, doi:10.1175/JHM-D-16-0119.1. 76.

GAUTAM, R. C. *et al.* Drought in India: its impact and mitigation strategies-a review. **Indian Journal of Agronomy**, v. 59, n. 2, p. 179-190, 2014.

GOERL, Roberto Fabris; KOBAYAMA, Masato. Redução dos desastres naturais: desafio dos geógrafos Natural disaster reduction: the challenge of geographers. **Ambiência**, v. 9, n. 1, p. 145-172, 2013.

GRIMM, A. M.; SABOIA, J. P. J. Interdecadal variability of the South American precipitation in the monsoon season. *Journal of Climate*, v. 28, n. 2, p. 755-775, 2015. DOI: 10.1175/JCLI-D-14-00046.1.

Haddad, M.S., 2011. Capacity choice and water management in hydroelectricity systems. *Energy Econ.* 33, 168–177. <https://doi.org/10.1016/j.eneco.2010.05.005>.

Han, Z.; Huang, S.; Huang, Q.; Leng, G.; Wang, H.; He, L.; Fang, W.; Li, P. Assessing GRACE-based terrestrial water storage anomalies dynamics at multi-timescales and their correlations with teleconnection factors in Yunnan Province, China. *J Hydrol.* 2019, 574, 836–850

Hartmann, H., Snow, J.A., Su, B., Jiang, T., 2016b. Seasonal predictions of precipitation in the Aksu-Tarim River basin for improved water resources management. *Glob. Planet. Chang.* 147, 86–96. <https://doi.org/10.1016/j.gloplacha.2016.10.018>.

HASTIE, Trevor; TIBSHIRANI, Robert. Exploring the nature of covariate effects in the proportional hazards model. **Biometrics**, p. 1005-1016, 1990.

HEALD, Colette L.; SPRACKLEN, Dominick V. Land use change impacts on air quality and climate. **Chemical reviews**, v. 115, n. 10, p. 4476-4496, 2015.

Huang, S.Z.; Huang, Q.; Chang, J.X.; Leng, G.Y. Linkages between hydrological drought, climate indices and human activities: A case study in the Columbia River Basin. *Int. J. Climatol.* 2016, 36, 280–290

Ionita M., Scholz P. e Chelcea S. 2016. Assessment of droughts in Romania using the Standardized Precipitation Index. *Natural Hazards*, 81(3), 1483–1498. doi: 10.1007/s11069-015-2141-8
Janizadeh, S., Avand, M., Jaafari, A., Phong, T.V., Bayat, M., Ahmadisharaf, E., Prakash, I., Pham, B.T., Lee, S., 2019. Prediction success of machine learning methods for flash flood susceptibility mapping in the tafresh watershed, Iran. *Sustainability* 11, 5426. <https://doi.org/10.3390/su11195426>.

KAYANO, M. T.; CAPISTRANO, V. B.; AMBRIZZI, T. On the Relations of South American Summer Precipitation Interannual Variability with Its Atlantic Ocean SSTs Related Modes and Influences. *Theoretical and Applied Climatology*, v. 124, p. 853-861, 2016. DOI: 10.1007/s00704-015-1459-x.

Khosravi, A., Nahavandi, S., Creighton, D., 2010. A prediction interval-based approach to determine optimal structures of neural network metamodels. *Expert Syst. Appl.* 37, 2377–2387. <https://doi.org/10.1016/j.eswa.2009.07.059>. Khosravi, K., Shahabi, H., Pham, B.T., Adamawoski, J., Shirzadi, A., Pradhan, B., Dou, J., Ly, H.-B., Gróf, G., Ho, H.L., 2019.

KHOSRAVI, Khabat *et al.* A comparative assessment of flood susceptibility modeling using multi-criteria decision-making analysis and machine learning methods. **Journal of Hydrology**, v. 573, p. 311-323, 2019.

KNIGHT, J. R.; FOLLAND, C. K.; SCOTT, P. A. A signature of persistent natural thermohaline circulation cycles in observed climate. *Geophysical Research Letters*, v. 33, n. 17, L17706, 2006. DOI: 10.1029/2006GL026242.

KUMAR, K. S.; RATHNAM, E. V.; SRIDHAR, V. Tracking seasonal and monthly drought with GRACE-based terrestrial water storage assessments over major river basins in South India. *Climate*, v. 9, n. 4, p. 56, 2021.

Kumar, K.S.; Rathnam, E.V.; Sridhar, V. Tracking seasonal and monthly drought with GRACE-based terrestrial water storage assessments over major river basins in South India. *Sci. Total Environ.* 2020, 763, 142994

LEJEUNE, Q.; EDOUARD, L. D.; BENOIT, P. G.; SONIA, I. S. Influence of Amazonian deforestation on the future evolution of regional surface fluxes, circulation, surface temperature and precipitation. *Climate Dynamics*, v. 44, p. 2769-2786, 2015. DOI: 10.1007/s00382-014-2203-8.

LEJEUNE, Q.; SÖRLIN, S.; SCHIETTECATTE, A.; DOMINGUES, P.; ABRAMS, J. A. Simulações do impacto do desmatamento na Amazônia com modelos climáticos regionais. *Climate Dynamics*, v. 45, n. 3-4, p. 987-1003, 2015.

- LEYDSON, J. B. Variabilidade Climática e Tendências de Precipitação no Nordeste Brasileiro. *Revista Brasileira de Meteorologia*, v. 35, n. 1, p. 111-121, 2020. DOI: 10.1590/0102-7786351015.
- LEYDSON, J. B. Variabilidade Climática e Tendências de Precipitação no Nordeste Brasileiro. *Revista Brasileira de Meteorologia*, v. 35, n. 1, p. 111-121, 2020. DOI: 10.1590/0102-7786351015.
- Liu, X.; Feng, X.; Ciais, P.; Fu, B. Widespread decline in terrestrial water storage and its link to teleconnections across Asia and eastern Europe. *Hydrol. Earth Syst. Sci.* 2020, 24, 3663–3676.
- MANATSA, D.; CHINGOMBE, W.; MATARIRA, C. H. The impact of the positive Indian Ocean dipole on Zimbabwe droughts Tropical climate is understood to be dominated by. *International Journal of Climatology*, v. 2029, n. March 2008, p. 2011–2029, 2008
- MANATSA, D.; CHINGOMBE, W.; MATARIRA, C. H. The impact of the positive Indian Ocean dipole on Zimbabwe droughts Tropical climate is understood to be dominated by. *International Journal of Climatology*, v. 2029, n. March 2008, p. 2011–2029, 2008
- MARENGO, J. A. Mudanças climáticas globais e seus efeitos sobre a biodiversidade: Caracterização do clima atual e definição das alterações climáticas para o território brasileiro ao longo do século XXI. São Paulo: Ministério do Meio Ambiente, 2004.
- Marengo, J. A., Nobre, C. A., Sampaio, G., Salazar, L. F. e Borma, L. S. (2011). Climate change in the Amazon Basin: Tipping points, changes in extremes, and impacts on natural and human systems. In *Tropical Rainforest Responses to Climatic Change* (pp. 259–283). Springer Praxis Books (ENVIRONSCI).
- MARENGO, J. A.; CHOU, S. C.; KAY, G.; ALVES, L. M.; PESQUERO, J. F.; SOARES, W. R.; SANTOS, D. C.; LYRA, A. A.; SUEIRO, G.; BETTS, R.; CHAGAS, D. J.; GOMES, J. L.; BUSSTAMANTE, J. F.; TAVARES, P. Development of regional future climate change scenarios in South America using the Eta CPTEC/HadCM3 climate change projections: climatology and regional analyses for the Amazon, São Francisco and the Paraná River basins. *Climate Dynamics*, v. 38, p. 1829-1848, 2012. DOI: 10.1007/s00382-011-1155-5.
- MARENGO, J. A.; HASTERNRATH, S. Case studies of extreme climatic events in the Amazon basin. *Journal of Climate*, v. 6, p. 617-627, 1993. DOI: 10.1175/1520-0442(1993)006<0617>2.0.CO;2.
- MARENGO, J. A.; TOMASELLA, J.; ALVES, L. M.; SOARES, W. R.; RODRIGUEZ, D. A. The drought of 2010 in the context of historical droughts in the Amazon region. *Geophysical Research Letters*, v. 38, L12703, 2011. DOI: 10.1029/2011GL047436.
- MARENGO, J. A.; TORRES, R. R.; ALVES, L. M. M. Drought in Northeast Brazil—Past, present, and future. *Applied Climate*, v. 129, p. 1189–1200, 2017. DOI: 10.1007/s00704-016-1840-8.
- Marengo, J.A.; Torres, R.R.; Alves, L.M.M. Drought in Northeast Brazil—Past, present, and future. *Appl. Clim.* 2017, 129, 1189–1200, doi:10.1007/s00704-016-1840-8.
- MARTHA JUNIOR, G. B.; CONTINI, E.; NAVARRO, Z. Caracterização do Amazônia Legal e macro-tendências do ambiente externo. Embrapa Estudos e Capacitação: Brasília, DF, 2011. 50 p.
- MCCULLOCH, C. E.; SEARLE, S. R. *Generalized, Linear, and Mixed Models*. 1.ed. United States of America: John Wiley e Sons, 2001. 358p
- MCCULLOCH, Charles E.; SEARLE, Shayle R. **Generalized, linear, and mixed models**. John Wiley e Sons, 2004.

Mehdizadeh, S.; Fathian, F.; Adamowski, J.D. Hybrid artificial intelligence-time series models for monthly streamflow modeling. *Appl. Soft Comput.* 2019, 80, 873–887, doi:10.1016/j.asoc.2019.03.046.

MEINKE, H.; STONE, R. C.; HAMMER, G. L.; HARRIS, G.; NELSON, R.; WHETTON, P. H.; WILKINSON, S.; WILSON, W.; CHAPMAN, S. C.; MERTON, A. A.; HOWDEN, S. M. Managing for climate variability in the grazing industry of northern Australia. *International Journal of Climatology*, v. 25, n. 14, p. 1947–1960, 2005. DOI: 10.1002/joc.1192.

MESTAS-NUÑEZ, A. M.; ENFIELD, D. B. Rotated global modes of non-ENSO sea surface temperature variability. *Journal of Climate*, v. 12, n. 10, p. 2734–2746, 1999. DOI: 10.1175/1520-0442(1999)012<2734>

Mouatadid, S., Raj, N., Deo, R.C., Adamowski, J.F., 2018. Input selection and data-driven model performance optimization to predict the standardized precipitation and evaporation index in a drought-prone region. *Atmos. Res.* 212, 130–149. <https://doi.org/10.1016/j.atmosres.2018.05.012>. MUHAMMAD, W., YANG, H., LEI, H., MUHAMMAD, A., YANG, D. (2018). Improving the regional applicability of satellite precipitation products by ensemble algorithm. *Remote Sens.* 10, 577. <https://doi.org/10.3390/rs10040577>

NAKAMURA, L. R. *et al.* Modelling location, scale and shape parameters of the birnbaum-saunders generalized t distribution. *Journal of Data Science*, v. 15, n. 2, p. 221–237, 2017. Ndehedehe, C.E.; Awange, J.L.; Kuhn, M.; Agutu, N.O.; Fukuda, Y. Climate teleconnections influence on West Africa's terrestrial water storage. *Hydrol. Process.* 2017, 31, 3206–3224.

NELDER, John Ashworth; WEDDERBURN, Robert WM. Generalized linear models. **Journal of the Royal Statistical Society: Series A (General)**, v. 135, n. 3, p. 370–384, 1972.

Ni, S.N.; Chen, J.L.; Wilson, C.R. Global Terrestrial Water Storage Changes and Connections to ENSO Events. *Surv. Geophys.* 2018, 39, 1–22.

NRAA. 2009. Drought Management Strategies – 2009. Draft paper of National Rainfed Area Authority, Ministry of Agriculture, Government of India, New Delhi, pp 70. PHAM, Binh Thai *et al.* Development of advanced artificial intelligence models for daily rainfall prediction. **Atmospheric Research**, v. 237, p. 104845, 2020.

Phillips, T.; Nerem, R.S.; Fox-Kemper, B.; Famiglietti, J.S.; Rajagopalan, B. The influence of ENSO on global terrestrial water storage using GRACE. *Geophys. Res. Lett.* 2012, 39, L16705.

PINKAYAN, S. S. Conditional Probability of Occurrence of Wet and Dry Years Over a Large Continental Area. *Journal of Applied Meteorology*, v. 5, n. 3, p. 375–385, 1966. DOI: 10.1175/1520-0450(1966)005<0375

RASHID, M. M.; AHMED, M. S.; HOSSAIN, M. I. Zero-inflated Poisson model for analyzing count data with excess zeros. *Journal of Statistical Computation and Simulation*, v. 86, n. 11, p. 2185–2198, 2016. DOI: 10.1080/00949655.2016.1140132.

RASHID, M.M., Beecham, M., Chowdhury, R.K., 2015. Assessment of trends in point rainfall using Continuous Wavelet Transforms. *Adv. Wat Res.*, 82, 1–15

Rashid, M.M.; Beechman, S. Development of a non-stationary Standardized Precipitation Index and its application to a South Australian climate. *Sci. Total Environ.* 2019, 657, 882–892, doi:10.1016/j.scitotenv.2018.12.052.

RIGBY, R. A.; STASINOPOULOS, D. M.; HELLER, G. Z.; DE BASTIANI, F. Distributions for modeling location, scale, and shape: Using GAMLSS in R. Chapman and Hall/CRC, 2019.

RIGBY, R. A.; STASINOPOULOS, D. M. Generalized additive models for location, scale and shape. *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, v.54, n.3, p.507–554, 2005.

RIGBY, R. A.; STASINOPOULOS, D. M. Generalized additive models for location, scale and shape (with discussion). *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, v. 54, n. 3, p. 507-554, 2005. DOI: 10.1111/j.1467-9876.2005.00510.x.

Rigby, R.A.; Stasinopoulos, D.M. Generalized additive models for location, scale and shape. *J. R. Stat. Soc. C Appl.* 2005, 54, 507–554, doi:10.1111/j.1467-9876.2005.00510.x.

Rio Grande do Sul – ABRH. 4ª Edição. Porto Alegre, 943 p

RODRIGUES, D. S. Análise de diagnóstico em modelos de regressão ZAGA e ZAIG. 2016. Dissertação (Mestrado em Estatística) - Universidade Estadual de Campinas, Campinas, 2016.

Roy, A.K. and Hirway, I. 2007. Multiple Impacts of Droughts and Assessment of Drought Policy in Major Drought Prone States in India. Project Report Submitted to the Planning Commission, Government of India, New Delhi.

ROY, Anil Kumar *et al.* Multiple impacts of droughts and assessment of drought policy in major drought prone states in India. **Gujarat, India: Centre for Development Alternatives**, 2007.

S. N. Wood. Generalized additive models. An introduction with R. Chapman e Hall, 2006a.

SAMRA, J. S. Review and analysis of drought monitoring, declaration and management in India. 2004.

SANTOS, Daris Correia dos. **Índices de extremos climáticos baseados na precipitação pluvial e influência de teleconexões no Nordeste do Brasil**. 2023. Dissertação (Mestrado em Engenharia Civil e Ambiental) - Programa de Pós-Graduação em Engenharia Civil e Ambiental, Universidade Federal da Paraíba, João Pessoa, 2023.

SATISH KUMAR, Kuruva *et al.* Monthly and seasonal drought characterization using grace-based groundwater drought index and its link to teleconnections across south indian river basins. **Climate**, v. 9, n. 4, p. 56, 2021.

Schneider U., Ziese M., Meyer-Christoffer A., Finger P., Rustemeier E. e Becker A. 2016. The new portfolio of global precipitation data products of the Global Precipitation Climatology Centre suitable to assess and quantify the global water cycle and resources. *Proceedings of the International Association of Hydrological Sciences*, 374, 29–34. doi: 10.5194/piahs-374-29-2016

SCHWARZ, G. *et al.* Estimating the dimension of a model. *The annals of statistics*, Institute of Mathematical Statistics, v. 6, n. 2, p. 461–464, 1978.

Serinaldi, F., Kilsby, C.G., 2012. A modular class of multisite monthly rainfall generators for water resource management and impact studies. *J. Hydrol.* 464–465, 528–540. <https://doi.org/10.1016/j.jhydrol.2012.07.043>.

Stasinopoulos D. M. and R. A. Rigby. Generalized additive models for location scale and shape (GAMLSS) in R. *Journal of Statistical Software*, 23(7):1–46, 2007.

STASINOPOULOS, M. D. *et al.* Flexible Regression and Smoothing: Using GAMLSS in R. [S.l.]: CRC Press, 2017.

STASINOPOULOS, M. D.; RIGBY, R. A.; BASTIANI, F. D. GAMLSS: a distributional regression approach. *Statistical Modelling*, v.18, n.3-4, p.248–273, 2018.

Stasinopoulos, M.D.; Rigby, R.A.; Heller, G.Z.; Voudouris, V.; De Bastiani, F. (Eds.) *Flexible Regression and Smoothing: Using GAMLSS in R*; Chapman and Hall/CRC The R Series: Boca Raton, FL, USA, 2017, doi:10.1201/b21973.

TAVARES, P.; GRIMM, A. M.; SABA, M. M. F. Interannual Variability of South American Precipitation and Its Links to Atmospheric Circulation and Sea Surface Temperature. *International Journal of Climatology*, v. 38, n. 8, p. 3231-3245, 2018. DOI: 10.1002/joc.5488.

Teodoro P.E., Oliveira-Júnior J.F., Cunha E.R., Correa C.C.G., Torres F.E., Bacani V.M., Gois G. e Ribeiro L.P. 2016. Cluster analysis applied to the spatial and temporal variability of monthly rainfall in Mato Grosso do Sul State, Brazil. *Meteorology and Atmospheric Physics*, 128(2), 197–209. doi: 10.1007/s00703-015-0408-y

Tien Bui, D., Khosravi, K., Shahabi, H., Daggupati, P., Adamowski, J.F., Melesse, A.M., Thai Pham, B., Pourghasemi, H.R., Mahmoudi, M., Bahrami, S., Pradhan, B., Shirzadi, A., Chapi, K., Lee, S., 2019. Flood spatial modeling in Northern Iran using remote sensing and GIS: a comparison between evidential belief functions and its ensemble with a multivariate logistic reg

Tien Bui, D., Pradhan, B., Nampak, H., Bui, Q.-T., Tran, Q.-A., Nguyen, Q.-P., 2016. Hybrid artificial intelligence approach based on neural fuzzy inference model and metaheuristic optimization for flood susceptibility modeling in a high-frequency tropical cyclone area using GIS. *J. Hydrol.* 540, 317–330. <https://doi.org/10.1016/j.jhydrol.2016.06.027>.

Trinh, T.A., 2018. The impact of climate change on agriculture: findings from households in Vietnam

TUCCI, C. E. M. (2009). *Hidrologia: ciência e aplicação*. Editora da Universidade Federal do

TURKMAN, A. A.; SILVA, G. L. *Modelos Lineares Generalizados - da teoria à prática*. Lisboa: [s.n.], 2000

Van Dijk A.I.J.M., Beck H.E., Crosbie R.S., Jeu R.A.M., Liu Y.Y., Podger G.M., Timbal B. e Viney N.R. 2013. The millennium drought in Southeast Australia (2001–2009): natural and human causes and implications for water resources, ecosystems, economy and society. *Water Resources Research*, 49(2), 1040–1057. doi: 10.1002/wrcr.20123

VOUDOURIS, V. *et al.* Modelling skewness and kurtosis with the bcpe density in gamlss. *Journal of Applied Statistics*, Taylor e Francis, v. 39, n. 6, p. 1279–1293, 2012.

Wang X., Shen H., Zhang W., Cao J., Qi Y., Chen G. e Li X. 2015. Spatial and temporal characteristics of droughts in the Northeast China Transect. *Natural Hazards*, 76(1), 601–614. doi: 10.1007/s11069-014-1507-7

Wang, F.; Wang, Z.; Yang, H.; Di, D.; Zhao, Y.; Liang, Q. Utilizing GRACE-based groundwater drought index for drought characterization and teleconnection factors analysis in the North China Plain. *J. Hydrol.* 2020, 585, 124849

XAVIER, T. M. B.; SANTOS, A. M.; CAVALCANTI, I. F. A. Previsão Climática Sazonal para o Nordeste do Brasil Utilizando a Técnica de Ensemble. *Revista Brasileira de Meteorologia*, v. 17, n. 2, p. 1-16, 2002.

YU, Manzhu; YANG, Chaowei; LI, Yun. Big data in natural disaster management: a review. *Geosciences*, v. 8, n. 5, p. 165, 2018.

Zhang A. e Jia G. 2013. Monitoring meteorological drought in semiarid regions using multisensor microwave remote sensing data. *Remote Sensing of Environment*, 134, 12–23. doi: 10.1016/j.rse.2013.02.023

ZHANG, R.; BUSALACCHI, A. J.; GU, D.; HALPERN, D.; MILLER, L. J.; MOORE, A.; SPERBER, K. R.; WALISER, D. E. Seasonal-to-interannual predictions with the Center for Ocean-Land-Atmosphere Studies coupled model system. *Bulletin of the American Meteorological Society*, v. 78, n. 4, p. 647-658, 1997. DOI: 10.1175/1520-0477(1997)078<0647>

Zhang, Z.; Chao, B.; Chen, J.; Wilson, C. Terrestrial water storage anomalies of Yangtze River basin droughts observed by GRACE and connections with ENSO. *Glob. Planet. Chang.* 2015, 126, 35–45

Zhao Q., Chen Q., Jiao M., Wu P., Gao X., Ma M. e Hong Y. 2018. The temporal-spatial characteristics of drought in the Loess Plateau using the remote-sensed TRMM precipitation data from 1998 to 2014. *Remote Sensing*, 10(6), 838. doi: 10.3390/rs10060838

Zhong R., Chen X., Lai C., Wang Z., Lian Y., Yu H. e Wu X. 2019. Drought monitoring utility of satellite-based precipitation products across mainland China. *Journal of Hydrology*, 568, 343–359. doi: 10.1016/j.jhydrol.2018.10.072

ZHOU, Lei *et al.* Emergency decision making for natural disasters: An overview. **International journal of disaster risk reduction**, v. 27, p. 567-576, 2018.

ZHOU, Lei *et al.* Emergency decision making for natural disasters: An overview. **International journal of disaster risk reduction**, v. 27, p. 567-576, 2018.

APÊNDICE A
CÓDIGOS NO R

`## SCRIPT DA DISSERTAÇÃO FAMÍLIA ZAGA`

`# Aplicando para a familia ZAGA e só para pixeis na AML`

```
# install.packages('hydroGOF')
library(hydroGOF)
library(gamlss)
library(gtools)
```

```
#### ESCOLHA A FAMÍLIA
familia <- 'ZAGA'
```

`# Configurar o diretório`

```
setwd("C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/PIXEIS_DENTRO_AML")
```

`# Lista todos os arquivos no diretório com extensão .xlsx`
`arquivos <- list.files(pattern = "\\\\.xlsx$")`

```
# Filtra apenas arquivos Excel
arquivos <- arquivos[grepl("\\\\.xlsx$", arquivos, ignore.case = TRUE)]
arquivos <- mixedsort(arquivos)
```

```
# Inicializar um dataframe para armazenar os resultados
resultados_df <- data.frame(Modelo = character(), R2 = numeric(), PBIAS = numeric(),
MAE = numeric(), stringsAsFactors = FALSE)
```

`# Loop sobre cada arquivo`

```
# Lista todos os arquivos no diretório com extensão .xlsx
arquivos <- list.files(pattern = "\\\\.xlsx$")
arquivos <- arquivos[grepl("\\\\.xlsx$", arquivos, ignore.case = TRUE)]
arquivos <- mixedsort(arquivos)
```

```
# Inicializar um dataframe para armazenar os resultados finais
resultados_df <- data.frame(Modelo = character(), R2 = numeric(), PBIAS = numeric(),
MAE = numeric(), stringsAsFactors = FALSE)
```

```
# Inicializar um dataframe para armazenar os call_summary de cada modelo
df_modelos <- data.frame(CallSummary = character(), stringsAsFactors = FALSE)
```

```
# Loop sobre cada arquivo
for (arquivo in arquivos) {
```

```

# Ler o arquivo
df_completo <- readxl::read_excel(path = arquivo)
df_completo$soi <- as.numeric(gsub(",", ".", df_completo$soi))
df_completo$a0[768] <- mean(df_completo$a0)

# Dividir o dataframe em conjunto de treino e teste
train_size <- nrow(df_completo) - 12
train_df <- df_completo[1:train_size, ]
test_df <- df_completo[(train_size + 1):nrow(df_completo), ]

if (sum(test_df$Precipitation) == 0) {
  r_squared <- 0
  pbias <- 0
  mae <- 0
  # Adicionar resultados ao dataframe resultados_df
  resultados_df <- rbind(resultados_df, data.frame(Modelo = arquivo, R2 = r_squared,
PBIAS = pbias, MAE = mae))

} else {
  # Ajustar o modelo
  fit0 <- gamlss(Precipitation ~ 1, family = familia, data = df_completo)
  fit1 <- stepAICAll.A(fit0, scope = list(
    lower = ~1,
    upper = ~ tmin + tmax + velo_vento + amo + nao + ao + pdo + nino1_2 + nino3_4 +
nino4 + nino3,
    sigma.scope = ~tmin + tmax + velo_vento + amo + nao + ao + pdo + nino1_2 +
nino3_4 + nino4 + nino3,
    nu.scope = ~tmin + tmax + velo_vento + amo + nao + ao + pdo + nino1_2 + nino3_4 +
nino4 + nino3,
    tau.scope = ~tmin + tmax + velo_vento + amo + nao + ao + pdo + nino1_2 + nino3_4
+ nino4 + nino3
  ), k = 2)

  # Extrair a chamada do resumo
  call_summary <- (fit1)$call

  # Criar o modelo PF1 com a chamada do resumo
  PF1 <- eval(call_summary)

  # Gerar wormplot (wp)
  wp_plot <- wp(fit1)

  # Nome do arquivo de saída para o plot

  plot_file <-
paste0("C:/Users/user/Desktop/Dissertacao_AM_LEGAL/SCRIPTS_DA DISSERTASSA
O/RESULTADOS_METRICAS_MODELO_WORMPLOTS/ZAGA/Model_pixel_",
tools::file_path_sans_ext(basename(arquivo)), "_plot.png")

```

```

# Construir o título dinâmico com base no nome do arquivo
plot_title <- paste("Pixel_", tools::file_path_sans_ext(basename(arquivo)), sep = "")

# Exportar o wormplot como imagem PNG
png(plot_file, width = 800, height = 800)
wp(fit1)
mtext(plot_title, side=3, line=0.5, cex=1, col="black", outer=F)
title(xlab = "Quantis dos desvios normalizados", line = 2)
title(ylab = "Desvio",line = 2)
dev.off()

# Obter o call_summary como string
call_summary <- capture.output(print(fit1$call))
call_summary <- paste(call_summary, collapse = " ")

# Adicionar ponto-e-vírgula ao final do call_summary
call_summary <- paste(call_summary, ";", sep = "")

# Salvar o call_summary no dataframe df_modelos
df_modelos <- rbind(df_modelos, data.frame(CallSummary = call_summary))

# Fazer previsões com o modelo PF1
predicted_values <- predict(PF1, newdata = test_df, type = "response")

# Valores reais para os últimos 12 meses
true_values <- test_df$Precipitation

# Agrupe os valores preditos e reais em trimestres
predicted_trimestrais <- sapply(split(predicted_values, gl(4, 3)), sum)
true_trimestrais <- sapply(split(true_values, gl(4, 3)), sum)

# Calcular o coeficiente de determinação R² usando br2
r_squared <- br2(sim = predicted_trimestrais, obs = true_trimestrais)
pbias <- mean((predicted_values - true_values) / true_values) * 100
mae <- mean(abs(predicted_values - true_values))

# Adicionar resultados ao dataframe resultados_df
resultados_df <- rbind(resultados_df, data.frame(Modelo = arquivo, R2 = r_squared,
PBIAS = pbias, MAE = mae))

  print(resultados_df)
}
}

# Exibir resultados finais
print(resultados_df)

# Exibir dataframe df_modelos com os call_summary

```

```

print(df_modelos)

write.csv(resultados_df, file =
'C:/Users/user/Desktop/Dissertacao_AM_LEGAL/SCRIPTS_DA DISSERTASSAO/RES
ULTADOS_METRICAS_MODELO_WORMPLOTS/ZAGA/resultados_df_familia_ZAG
A.csv', row.names = TRUE)
write.csv(df_modelos, file =
'C:/Users/user/Desktop/Dissertacao_AM_LEGAL/SCRIPTS_DA DISSERTASSAO/MO
DELOS_FAMILIA_CSV/ZAGA/modelo_familia_ZAGA.csv', row.names = TRUE)

# Ler uma planilha específica do arquivo Excel
coord <-
readxl::read_excel('C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/coordenadas_de
ntro_aml.xlsx')
coord <- coord[,c(1,2)]

#coord <- coord[c(seq(1,3)),c(2,3)]

#coord <- coord[c(seq(1,nrow(coord)-2)),]
coord <- cbind(coord,resultados_df$R2)

colnames(coord)[colnames(coord) == "resultados_df$R2"] <- "R2"

df <- coord
library(ggplot2)
# Instale e carregue os pacotes necessários
#install.packages("ggplot2")
library(ggplot2)
library(sf)
library(dplyr)

# Carregue o shapefile (substitua com o caminho correto para o seu arquivo)
shapefile_path <-
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/L
imites_Amazonia_Legal_2022.shp"
sertao_sf <- st_read(shapefile_path)

# Defina o CRS para o shapefile (substitua "EPSG:4326" pelo CRS apropriado)
sertao_sf <- st_set_crs(sertao_sf, "EPSG:4326")

sertao_sf <- st_transform(sertao_sf, "EPSG:4326")

# Crie a grade de resolução 0.5 grau

```

```

grid <- st_make_grid(sertao_sf, cellsize = c(1, 1), what = "polygons")

# Seu dataframe de dados (substitua com seu próprio dataframe)
df <- coord
# Converta o dataframe em um objeto sf e defina o CRS (substitua "EPSG:4326" pelo CRS
apropriado)
df_sf <- st_as_sf(df, coords = c("Longitude", "Latitude"), crs = "EPSG:4326")

# Transforme o objeto sf se necessário (substitua "EPSG:4326" pelo CRS desejado)
df_sf <- st_transform(df_sf, crs = "EPSG:4326")

# Filtrar os dados para excluir os pixels com valor zero
df_sf_filtered <- df_sf[df_sf$R2 != 0, ]

# Crie o gráfico usando ggplot2
#meu_grafico <- ggplot() +
# geom_sf(data = sertao_sf, color = "darkblue", size = 1.5) +
# geom_tile(data = df_sf_filtered, aes(x = st_coordinates(geometry)[, 1], y =
st_coordinates(geometry)[, 2], fill = R2), color = "black", alpha = 0.2) +
# geom_sf(data = grid, color = NA, size = 0.2, alpha = 0.5) +
# scale_fill_gradient(low = "red", high = "green", limits = c(0.9, 1)) + # Ajuste os limites
da escala de cores
# theme_minimal() +
# ggtitle("Valores de R² dos modelos no Sertão da Paraíba") +
# labs(x = NULL, y = NULL)
#
#
#meu_grafico

df$R2[df$R2 < 0] <- 0

#tentativa 2:

meu_grafico <- ggplot() +
  geom_sf(data = sertao_sf, color = "darkblue", size = 1.5) +
  geom_tile(data = df_sf_filtered, aes(x = st_coordinates(geometry)[, 1], y =
st_coordinates(geometry)[, 2], fill = R2), color = "black") +
  geom_sf(data = grid, color = NA, size = 0.2, alpha = 0.5) +
  scale_fill_gradient(low = "red", high = "green", limits = c(0, 1)) + # Ajuste os limites da
escala de cores
  theme_minimal() +
  ggtitle("Valores de R² para a AML dos modelos com distribuição GG") +
  labs(x = NULL, y = NULL)

```

meu_grafico

#tentativa 3:

Carregar o shapefile

shapefile_path <-

"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/Limites_Amazonia_Legal_2022.shp"

sertao_sf <- st_read(shapefile_path)

Plot

ggplot() +

geom_sf(data = sertao_sf, color = "black", fill = NA, size = 2) + # Adiciona o shapefile

geom_point(data = df, aes(x = Longitude, y = Latitude, color = R2), size = 3) +

scale_color_gradient(low = "red", high = "green") +

labs(title = "Mapa de R2 com Limites da Amazônia Legal", x = "Longitude", y =

"Latitude") +

theme_minimal()

so as quadriculas

Carregar o shapefile

shapefile_path <-

"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/Limites_Amazonia_Legal_2022.shp"

sertao_sf <- st_read(shapefile_path)

Criar uma grade de resolução 0.5 grau

grid <- st_make_grid(sertao_sf, cellsize = c(1, 1), what = "polygons")

Plot coerente

ggplot() +

geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = R2), color = "black", alpha = 0.5) +

geom_sf(data = sertao_sf, color = "black", fill = NA, size = 3.0) + # Adiciona o shapefile+

scale_fill_gradient(low = "red", high = "green") +

labs(title = "Mapa de R2 com Limites da Amazônia Legal", x = "Longitude", y =

"Latitude") +

theme_minimal()

Exibir o gráfico

#(meu_grafico)

GRAFICO CERTO 01


```
# Plot com quadriculas de 1 grau
ggplot() +
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = R2), color = "black", width =
1, height = 1) +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 2) + # Adiciona o shapefile

  scale_fill_gradient(low = "red", high = "green") +
  labs(title = "Mapa de R2 com Limites da Amazônia Legal", x = "Longitude", y =
"Latitude") +
  theme_minimal()
```

#gGRAFICO CERTO 02

```
# Plot com quadriculas de 1 grau
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 8) + # Adiciona o shapefile
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = R2), color = "black", width =
1, height = 1,alpha = 0.6) +
  scale_fill_gradient(low = "red", high = "green") +
  labs(title = "Mapa de R2 com Limites da Amazônia Legal", x = "Longitude", y =
"Latitude") +
  theme_minimal()
```

#GRAFICO CERTO 03

```
# Plot com quadriculas
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 2) + # Adiciona o shapefile
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = R2), color = "black", width =
0.5, height = 0.5) +
  scale_fill_gradient(low = "red", high = "green") +
  labs(title = "Mapa de R2 com Limites da Amazônia Legal", x = "Longitude", y =
"Latitude") +
  theme_minimal()
```

Plot de quadriculas com r2 maiores que 0.8

```
df_maior_80 <-df[df$R2>0.80,]
# Plot com quadriculas de 1 grau
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 8) + # Adiciona o shapefile
  geom_tile(data = df_maior_80, aes(x = Longitude, y = Latitude, fill = R2), color =
"black", width = 1, height = 1,alpha = 0.6) +
  scale_fill_gradient(low = "yellow", high = "Darkgreen") +
```

```
labs(title = "Mapa de R2 com Limites da Amazônia Legal", x = "Longitude", y =
"Latitude") +
theme_minimal()
```

```
df_maior_80 <-df[df$R2>0.90,]
# Plot com quadriculas de 1 grau
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 8) + # Adiciona o shapefile
  geom_tile(data = df_maior_80, aes(x = Longitude, y = Latitude, fill = R2), color =
"black", width = 1, height = 1,alpha = 0.6) +
  scale_fill_gradient(low = "green", high = "Darkgreen") +
  labs(title = "Mapa de Pixels com R2 superior a 0.90", x = "Longitude", y = "Latitude") +
  theme_minimal()
```

Como o r2 está sendo calculado? Demonstração:

```
# Calcular a média dos valores previstos e observados
mean_predicted <- mean(predicted_trimestrais)
mean_true <- mean(true_trimestrais)
```

```
# Calcular a soma dos quadrados das diferenças entre os valores previstos e observados
ss_total <- sum((true_trimestrais - mean_true)^2)
```

```
# Calcular a soma dos quadrados das diferenças entre os valores previstos e a média
prevista
ss_residual <- sum((true_trimestrais - predicted_trimestrais)^2)
```

```
# Calcular R2
r_squared <- 1 - (ss_residual / ss_total)
```

```
# Se o valor de R2 for menor que zero, torná-lo zero
r_squared <- max(0, r_squared)
```

```
# Imprimir o valor de R2
print(paste("O coeficiente de determinação R2 é:", r_squared))
```

CÓDIGO DA ANÁLISE QUALITATIVA

SCRIPT DA DISSERTAÇÃO FAMÍLIA ZAGA

Aplicando para a família ZAGA e só para pixeis na AML

```
# install.packages('hydroGOF')
library(hydroGOF)
```

```

library(gamlss)
library(gtools)
library(readxl)
library(dplyr)
library(irr)
library(stats)

familia <- 'ZAGA'

# Configurar o diretório
setwd("C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/PIXEIS_DENTRO_AML")

# Diretório para salvar os CSVs
output_dir <-
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/Analise_qualitativa/dados_cada_model
o_dentro_da_aml"
dir.create(output_dir, recursive = TRUE, showWarnings = FALSE)

# Lista todos os arquivos no diretório com extensão .xlsx
arquivos <- list.files(pattern = "\\\\.xlsx$")
arquivos <- mixedsort(arquivos)

# Inicializar um dataframe para armazenar os resultados
resultados_df <- data.frame(Modelo = character(), R2 = numeric(), PBIAS = numeric(),
MAE = numeric(),
                             Kappa = numeric(), Kendall = numeric(), Predicted_Class = character(),
True_Class = character(), stringsAsFactors = FALSE)

# Função para classificar os valores conforme a tabela fornecida
classificar <- function(valores) {
  quantis <- quantile(valores, probs = c(0.05, 0.15, 0.35, 0.65, 0.85, 0.95))
  sapply(valores, function(x) {
    if (x < quantis[1]) {
      "Extremamente Seco"
    } else if (x < quantis[2]) {
      "Muito Seco"
    } else if (x < quantis[3]) {
      "Seco"
    } else if (x < quantis[4]) {
      "Normalidade"
    } else if (x < quantis[5]) {
      "Chuvoso"
    } else if (x < quantis[6]) {
      "Muito Chuvoso"
    } else {
      "Extremamente Chuvoso"
    }
  })
}

```

```

# Loop sobre cada arquivo
for (arquivo in arquivos) {

  # Ler o arquivo
  df_completo <- readxl::read_excel(path = arquivo)
  df_completo$soi <- as.numeric(gsub(",", ".", df_completo$soi))
  df_completo$ao[768] <- mean(df_completo$ao)

  # Dividir o dataframe em conjunto de treino e teste
  train_size <- nrow(df_completo) - 12
  train_df <- df_completo[1:train_size, ]
  test_df <- df_completo[(train_size + 1):nrow(df_completo), ]

  if (sum(test_df$Precipitation) == 0) {
    r_squared <- 0
    pbias <- 0
    mae <- 0
    kappa <- NA
    kendall <- NA
    predicted_classes <- rep(NA, 4)
    true_classes <- rep(NA, 4)
    # Adicionar resultados ao dataframe resultados_df
    resultados_df <- rbind(resultados_df, data.frame(Modelo = arquivo, R2 = r_squared,
PBIAS = pbias, MAE = mae,
                                     Kappa = kappa, Kendall = kendall, Predicted_Class =
predicted_classes, True_Class = true_classes))

  } else {
    # Ajustar o modelo
    fit0 <- gamlss(Precipitation ~ 1, family = familia, data = df_completo)
    fit1 <- stepAICAll.A(fit0, scope = list(
      lower = ~1,
      upper = ~ tmin + tmax + velo_vento + amo + nao + ao + pdo + nino1_2 + nino3_4 +
nino4 + nino3,
      sigma.scope = ~tmin + tmax + velo_vento + amo + nao + ao + pdo + nino1_2 +
nino3_4 + nino4 + nino3,
      nu.scope = ~tmin + tmax + velo_vento + amo + nao + ao + pdo + nino1_2 + nino3_4 +
nino4 + nino3,
      tau.scope = ~tmin + tmax + velo_vento + amo + nao + ao + pdo + nino1_2 + nino3_4
+ nino4 + nino3
    ), k = 2)

    # Fazer previsões com o modelo PF1
    predicted_values <- predict(fit1, newdata = test_df, type = "response")

    # Valores reais para os últimos 12 meses
    true_values <- test_df$Precipitation

    # Agrupar os valores preditos e reais em trimestres
    predicted_trimestrais <- sapply(split(predicted_values, gl(4, 3)), sum)
  }
}

```

```

true_trimestrais <- sapply(split(true_values, gl(4, 3)), sum)

# Calcular o coeficiente de determinação R² usando br2
r_squared <- br2(sim = predicted_trimestrais, obs = true_trimestrais)
pbias <- mean((predicted_values - true_values) / true_values) * 100
mae <- mean(abs(predicted_values - true_values))

# Classificar os valores trimestrais conforme a tabela fornecida
predicted_classes <- classificar(predicted_trimestrais)
true_classes <- classificar(true_trimestrais)

# Calcular o coeficiente kappa
kappa <- kappa2(data.frame(predicted_classes, true_classes))$value

# Calcular a correlação de Kendall
kendall <- cor(predicted_trimestrais, true_trimestrais, method = "kendall")

# Adicionar resultados ao dataframe resultados_df
resultados_df <- rbind(resultados_df, data.frame(Modelo = arquivo, R2 = r_squared,
PBIAS = pbias, MAE = mae,
                                     Kappa = kappa, Kendall = kendall, Predicted_Class =
paste(predicted_classes, collapse = "; "),
                                     True_Class = paste(true_classes, collapse = "; ")))

# Salvar os valores trimestrais em um CSV
trimestrais_df <- data.frame(Modelo = arquivo,
                             Predicted_Trim_1 = predicted_trimestrais[1], Predicted_Trim_2 =
predicted_trimestrais[2],
                             Predicted_Trim_3 = predicted_trimestrais[3], Predicted_Trim_4 =
predicted_trimestrais[4],
                             True_Trim_1 = true_trimestrais[1], True_Trim_2 =
true_trimestrais[2],
                             True_Trim_3 = true_trimestrais[3], True_Trim_4 =
true_trimestrais[4])

write.csv(trimestrais_df, file.path(output_dir, paste0("trimestrais_",
tools::file_path_sans_ext(basename(arquivo)), ".csv")), row.names = FALSE)

print(resultados_df)
}
}

# Salvar o dataframe final em um CSV
write.csv(resultados_df, file.path(output_dir, "analise_qualitativa_resultados_final.csv"),
row.names = FALSE)

# Visualizar o dataframe final
print(resultados_df)

```

```

# Deu um pau no pixel 85, vamos retirar os valores duplicados presentes nas linhas 7,8,9
resultados_df <- resultados_df[-c(7, 8, 9), ]

#####-----

# Buscar plotar kendall

# Ler uma planilha específica do arquivo Excel
coord <-
readxl::read_excel('C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/coordenadas_de
ntro_aml.xlsx')
coord <- coord[,c(1,2)]

#coord <- coord[c(seq(1,3)),c(2,3)]

#coord <- coord[c(seq(1,nrow(coord)-2)),]
coord_kappa <- cbind(coord,resultados_df$Kappa)

colnames(coord_kappa)[colnames(coord_kappa) == "resultados_df$Kappa"] <- "kappa"

df <- coord_kappa

df$kappa[df$kappa < 0] <- 0

library(ggplot2)
# Instale e carregue os pacotes necessários
#install.packages("ggplot2")
library(ggplot2)
library(sf)
library(dplyr)

# Carregue o shapefile (substitua com o caminho correto para o seu arquivo)
shapefile_path <-
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/L
imites_Amazonia_Legal_2022.shp"
sertao_sf <- st_read(shapefile_path)

# Defina o CRS para o shapefile (substitua "EPSG:4326" pelo CRS apropriado)
sertao_sf <- st_set_crs(sertao_sf, "EPSG:4326")

sertao_sf <- st_transform(sertao_sf, "EPSG:4326")

# Crie a grade de resolução 0.5 grau

```

```

grid <- st_make_grid(sertao_sf, cellsize = c(1, 1), what = "polygons")

# Seu dataframe de dados (substitua com seu próprio dataframe)
df <- coord
# Converta o dataframe em um objeto sf e defina o CRS (substitua "EPSG:4326" pelo CRS
apropriado)
df_sf <- st_as_sf(df, coords = c("Longitude", "Latitude"), crs = "EPSG:4326")

# Transforme o objeto sf se necessário (substitua "EPSG:4326" pelo CRS desejado)
df_sf <- st_transform(df_sf, crs = "EPSG:4326")

# Filtrar os dados para excluir os pixels com valor zero
df_sf_filtered <- df_sf[df_sf$kappa != 0, ]

```

GRAFICO CERTO 01

```

# Plot com quadriculas de 1 grau
ggplot() +
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = kappa), color = "black", width
= 1, height = 1) +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 2) + # Adiciona o shapefile

  scale_fill_gradient(low = "red", high = "green") +
  labs(title = "Mapa dos valores de Kappa ", x = "Longitude", y = "Latitude") +
  theme_minimal()

```

#gGRAFICO CERTO 02

```

# Plot com quadriculas de 1 grau
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 8) + # Adiciona o shapefile
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = kappa), color = "black", width
= 1, height = 1,alpha = 0.6) +
  scale_fill_gradient(low = "red", high = "green") +
  labs(title = "Mapa dos valores de Kappa", x = "Longitude", y = "Latitude") +
  theme_minimal()

```

#GRAFICO CERTO 03

```

# Plot com quadriculas
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 2) + # Adiciona o shapefile
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = R2), color = "black", width =
0.5, height = 0.5) +

```

```

scale_fill_gradient(low = "red", high = "green") +
labs(title = "Mapa de R2 com Limites da Amazônia Legal", x = "Longitude", y =
"Latitude") +
theme_minimal()

```

Plot de quadriculas com r2 maiores que 0.8

```

df_maior_80 <-df[df$kappa==1,]
# Plot com quadriculas de 1 grau
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 8) + # Adiciona o shapefile
  geom_tile(data = df_maior_80, aes(x = Longitude, y = Latitude, fill = kappa), color =
"black", width = 1, height = 1,alpha = 0.6) +
  scale_fill_gradient(low = "yellow", high = "Darkgreen") +
  labs(title = "Mapa com Kappa = 1", x = "Longitude", y = "Latitude") +
  theme_minimal()

```

#####-----

```

# Ler uma planilha específica do arquivo Excel
coord <-
readxl::read_excel('C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/coordenadas_de
ntro_aml.xlsx')
coord <- coord[,c(1,2)]

```

```
#coord <- coord[c(seq(1,3)),c(2,3)]
```

```
#coord <- coord[c(seq(1,nrow(coord)-2)),]
coord_kendall <- cbind(coord,resultados_df$Kendall)
```

```
colnames(coord_kendall)[colnames(coord_kendall) == "resultados_df$Kendall"] <-
"kendall"
```

```
df <- coord_kendall
```

```
df <- drop_na(df)
```

```
library(ggplot2)
```



```

# Instale e carregue os pacotes necessários
#install.packages("ggplot2")
library(ggplot2)
library(sf)
library(dplyr)

# Carregue o shapefile (substitua com o caminho correto para o seu arquivo)
shapefile_path <-
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/L
imites_Amazonia_Legal_2022.shp"
sertao_sf <- st_read(shapefile_path)

# Defina o CRS para o shapefile (substitua "EPSG:4326" pelo CRS apropriado)
sertao_sf <- st_set_crs(sertao_sf, "EPSG:4326")

sertao_sf <- st_transform(sertao_sf, "EPSG:4326")

# Crie a grade de resolução 0.5 grau
grid <- st_make_grid(sertao_sf, cellsize = c(1, 1), what = "polygons")

# Converta o dataframe em um objeto sf e defina o CRS (substitua "EPSG:4326" pelo CRS
apropriado)
df_sf <- st_as_sf(df, coords = c("Longitude", "Latitude"), crs = "EPSG:4326")

# Transforme o objeto sf se necessário (substitua "EPSG:4326" pelo CRS desejado)
df_sf <- st_transform(df_sf, crs = "EPSG:4326")

# GRAFICO CERTO 01

# Plot com quadriculas de 1 grau
ggplot() +
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = kendall), color = "black",
width = 1, height = 1) +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 2) + # Adiciona o shapefile

  scale_fill_gradient(low = "red", high = "green") +
  labs(title = "Mapa dos valores de Kendall ", x = "Longitude", y = "Latitude") +
  theme_minimal()

#gGRAFICO CERTO 02

# Plot com quadriculas de 1 grau
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 8) + # Adiciona o shapefile
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = kendall), color = "black",
width = 1, height = 1,alpha = 0.6) +

```

```
scale_fill_gradient(low = "red", high = "green") +
labs(title = "Mapa dos valores de Kappa", x = "Longitude", y = "Latitude") +
theme_minimal()
```

#GRAFICO CERTO 03

```
# Plot com quadriculas
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 2) + # Adiciona o shapefile
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = kendall), color = "black",
width = 0.5, height = 0.5) +
  scale_fill_gradient(low = "red", high = "green") +
  labs(title = "Mapa de R2 com Limites da Amazônia Legal", x = "Longitude", y =
"Latitude") +
  theme_minimal()
```

Plot de quadriculas com r2 maiores que 0.8

```
df_maior_80 <-df[df$kendall==1,]
# Plot com quadriculas de 1 grau
ggplot() +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 8) + # Adiciona o shapefile
  geom_tile(data = df_maior_80, aes(x = Longitude, y = Latitude, fill = kendall), color =
"black", width = 1, height = 1,alpha = 0.6) +
  scale_fill_gradient(low = "yellow", high = "Darkgreen") +
  labs(title = "Mapa com Kendall = 1", x = "Longitude", y = "Latitude") +
  theme_minimal()
```

#####-----

```
# Ler uma planilha específica do arquivo Excel
coord <-
readxl::read_excel('C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/coordenadas_de
ntro_aml.xlsx')
coord <- coord[,c(1,2)]
```

```
#coord <- coord[c(seq(1,3)),c(2,3)]
```

```
#coord <- coord[c(seq(1,nrow(coord)-2)),]
coord <- cbind(coord,resultados_df)
```

```
colnames(coord)[colnames(coord) == "resultados_df$R2"] <- "R2"
```

```

df <- coord
library(ggplot2)
# Instale e carregue os pacotes necessários
#install.packages("ggplot2")
library(ggplot2)
library(sf)
library(dplyr)

# Carregue o shapefile (substitua com o caminho correto para o seu arquivo)
shapefile_path <-
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/L
imites_Amazonia_Legal_2022.shp"
sertao_sf <- st_read(shapefile_path)

# Defina o CRS para o shapefile (substitua "EPSG:4326" pelo CRS apropriado)
sertao_sf <- st_set_crs(sertao_sf, "EPSG:4326")

sertao_sf <- st_transform(sertao_sf, "EPSG:4326")

# Crie a grade de resolução 0.5 grau
grid <- st_make_grid(sertao_sf, cellsize = c(1, 1), what = "polygons")

# Seu dataframe de dados (substitua com seu próprio dataframe)
df <- coord
# Converta o dataframe em um objeto sf e defina o CRS (substitua "EPSG:4326" pelo CRS
apropriado)
df_sf <- st_as_sf(df, coords = c("Longitude", "Latitude"), crs = "EPSG:4326")

# Transforme o objeto sf se necessário (substitua "EPSG:4326" pelo CRS desejado)
df_sf <- st_transform(df_sf, crs = "EPSG:4326")

# Filtrar os dados para excluir os pixels com valor zero
df_sf_filtered <- df_sf[df_sf$R2 != 0, ]

# Plot com quadriculas de 1 grau
ggplot() +
  geom_tile(data = df, aes(x = Longitude, y = Latitude, fill = R2), color = "black", width =
1, height = 1) +
  geom_sf(data = sertao_sf, color = "black", fill = NA, size = 2) + # Adiciona o shapefile

  scale_fill_gradient(low = "red", high = "green") +
  labs(title = "Mapa de R2 com Limites da Amazônia Legal", x = "Longitude", y =
"Latitude") +
  theme_minimal()

```

```
## tentativa do pinoquio
```

```
# Carregar pacotes necessários
```

```
library(ggplot2)
```

```
library(sf)
```

```
library(dplyr)
```

```
library(tidyr)
```

```
library(readxl)
```

```
# Ler as coordenadas e os resultados
```

```
coord <-
```

```
readxl::read_excel('C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/coordenadas_de  
ntro_aml.xlsx')
```

```
coord <- coord[, c(1, 2)]
```

```
coord <- cbind(coord, resultados_df)
```

```
colnames(coord)[colnames(coord) == "resultados_df$R2"] <- "R2"
```

```
# Carregar o shapefile
```

```
shapefile_path <-
```

```
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/L  
imites_Amazonia_Legal_2022.shp"
```

```
sertao_sf <- st_read(shapefile_path)
```

```
# Definir e transformar o CRS
```

```
sertao_sf <- st_set_crs(sertao_sf, "EPSG:4326")
```

```
sertao_sf <- st_transform(sertao_sf, crs = "EPSG:4326")
```

```
# Separar os valores de cada trimestre em colunas separadas
```

```
coord <- coord %>%
```

```
  separate(Predicted_Class, into = paste0("Predicted_Trim_", 1:4), sep = ";") %>%
```

```
  separate(True_Class, into = paste0("True_Trim_", 1:4), sep = ";")
```

```
# Converter o dataframe em um objeto sf e definir o CRS
```

```
df_sf <- st_as_sf(coord, coords = c("Longitude", "Latitude"), crs = "EPSG:4326")
```

```
df_sf <- st_transform(df_sf, crs = "EPSG:4326")
```

```
# Filtrar os dados para excluir os pixels com valor zero
```

```
df_sf_filtered <- df_sf[df_sf$R2 != 0, ]
```

```
# Ajustar a posição e o tamanho do texto da legenda
```

```
legend_theme <- theme(
```

```

legend.position = "bottom",
legend.title = element_text(size = 10),
legend.text = element_text(size = 8)
)

# Função para criar mapas individuais para cada trimestre
plot_map <- function(data, column, title) {
  ggplot() +
    geom_sf(data = sertao_sf, fill = "white", color = "black") +
    geom_sf(data = data, aes_string(color = column)) +
    theme_minimal() +
    labs(title = title, color = "Classificação") +
    legend_theme
}

# Criar os mapas para cada trimestre
mapas_predicted <- lapply(1:4, function(i) {
  plot_map(df_sf_filtered, paste0("Predicted_Trim_", i), paste("Predicted Class - Trim", i))
})

mapas_true <- lapply(1:4, function(i) {
  plot_map(df_sf_filtered, paste0("True_Trim_", i), paste("True Class - Trim", i))
})

# Mostrar os gráficos individualmente
lapply(mapas_predicted, print)
lapply(mapas_true, print)

# Salvar os gráficos como imagens separadas
for (i in 1:4) {

  ggsave(paste0("C:/Users/user/Desktop/Dissertacao_AM_LEGAL/Analise_qualitativa/mapa_predicted_trim_", i, ".png"), plot = mapas_predicted[[i]], width = 7, height = 7, units = "in")

  ggsave(paste0("C:/Users/user/Desktop/Dissertacao_AM_LEGAL/Analise_qualitativa/mapa_true_trim_", i, ".png"), plot = mapas_true[[i]], width = 7, height = 7, units = "in")
}

# plotando os trimestres um ao lado do outro:

# Carregar pacotes necessários
library(ggplot2)

```

```

library(sf)
library(dplyr)
library(tidyr)
library(readxl)
library(patchwork)

# Ler as coordenadas e os resultados
coord <-
readxl::read_excel('C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/coordenadas_de
ntro_aml.xlsx')
coord <- coord[, c(1, 2)]
coord <- cbind(coord, resultados_df)
colnames(coord)[colnames(coord) == "resultados_df$R2"] <- "R2"

# Carregar o shapefile
shapefile_path <-
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/L
imites_Amazonia_Legal_2022.shp"
sertao_sf <- st_read(shapefile_path)

# Definir e transformar o CRS
sertao_sf <- st_set_crs(sertao_sf, "EPSG:4326")
sertao_sf <- st_transform(sertao_sf, crs = "EPSG:4326")

# Separar os valores de cada trimestre em colunas separadas
coord <- coord %>%
  separate(Predicted_Class, into = paste0("Predicted_Trim_", 1:4), sep = ";") %>%
  separate(True_Class, into = paste0("True_Trim_", 1:4), sep = ";")

# Converter o dataframe em um objeto sf e definir o CRS
df_sf <- st_as_sf(coord, coords = c("Longitude", "Latitude"), crs = "EPSG:4326")
df_sf <- st_transform(df_sf, crs = "EPSG:4326")

# Filtrar os dados para excluir os pixels com valor zero
df_sf_filtered <- df_sf[df_sf$R2 != 0, ]

# Criar a grade de resolução 0.5 grau
grid <- st_make_grid(sertao_sf, cellsize = c(0.5, 0.5), what = "polygons")

# Ajustar a posição e o tamanho do texto da legenda
legend_theme <- theme(
  legend.position = "bottom",
  legend.title = element_text(size = 10),
  legend.text = element_text(size = 8)
)

# Função para criar mapas individuais para cada trimestre
plot_map <- function(data, column, title) {
  ggplot() +
    geom_sf(data = sertao_sf, fill = "white", color = "black") +

```

```

    geom_sf(data = st_as_sf(st_intersection(data, st_as_sf(grid))), aes_string(fill = column))
+
    theme_minimal() +
    labs(title = title, fill = "Classificação") +
    legend_theme
}

# Criar os mapas para cada trimestre
mapas <- lapply(1:4, function(i) {
  predicted_map <- plot_map(df_sf_filtered, paste0("Predicted_Trim_", i), paste("Predicted
Class - Trim", i))
  true_map <- plot_map(df_sf_filtered, paste0("True_Trim_", i), paste("True Class - Trim",
i))
  combined_map <- predicted_map + true_map + plot_layout(ncol = 2)
  return(combined_map)
})

# Mostrar os gráficos
lapply(mapas, print)

# Salvar os gráficos como imagens separadas
for (i in 1:4) {

  ggsave(paste0("C:/Users/user/Desktop/Dissertacao_AM_LEGAL/Analise_qualitativa/com
paracao/mapa_comparacao_trim_", i, ".png"), plot = mapas[[i]], width = 14, height = 7,
units = "in")
}

# Carregar pacotes necessários
library(ggplot2)
library(sf)
library(dplyr)
library(tidyr)
library(readxl)
library(gridExtra)

# Verificar se há algum dispositivo gráfico aberto e fechá-lo
if (!is.null(dev.list())) dev.off()

# Ler as coordenadas e os resultados
coord <-
readxl::read_excel('C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/coordenadas_de
ntro_aml.xlsx')
coord <- coord[, c(1, 2)]
coord <- cbind(coord, resultados_df)
colnames(coord)[colnames(coord) == "resultados_df$R2"] <- "R2"

# Carregar o shapefile

```

```

shapefile_path <-
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/L
imites_Amazonia_Legal_2022.shp"
sertao_sf <- st_read(shapefile_path)

# Definir e transformar o CRS
sertao_sf <- st_set_crs(sertao_sf, "EPSG:4326")
sertao_sf <- st_transform(sertao_sf, crs = "EPSG:4326")

# Separar os valores de cada trimestre em colunas separadas
coord <- coord %>%
  separate(Predicted_Class, into = paste0("Predicted_Trim_", 1:4), sep = ";") %>%
  separate(True_Class, into = paste0("True_Trim_", 1:4), sep = ";")

# Converter o dataframe em um objeto sf e definir o CRS
df_sf <- st_as_sf(coord, coords = c("Longitude", "Latitude"), crs = "EPSG:4326")
df_sf <- st_transform(df_sf, crs = "EPSG:4326")

# Filtrar os dados para excluir os pixels com valor zero
df_sf_filtered <- df_sf[df_sf$R2 != 0, ]

# Criar a grade de resolução 0.5 grau
grid <- st_make_grid(sertao_sf, cellsize = c(0.5, 0.5), what = "polygons")

# Ajustar a posição e o tamanho do texto da legenda
legend_theme <- theme(
  legend.position = "bottom",
  legend.title = element_text(size = 10),
  legend.text = element_text(size = 8)
)

# Função para criar mapas individuais para cada trimestre
plot_map <- function(data, column, title) {
  ggplot() +
    geom_sf(data = sertao_sf, fill = "white", color = "black") +
    geom_sf(data = st_as_sf(st_intersection(data, st_as_sf(grid))), aes_string(fill = column))
  +
    theme_minimal() +
    labs(title = title, fill = "Classificação") +
    legend_theme
}

# Criar os mapas para cada trimestre
mapas <- lapply(1:4, function(i) {
  predicted_map <- plot_map(df_sf_filtered, paste0("Predicted_Trim_", i), paste("Predicted
Class - Trim", i))
  true_map <- plot_map(df_sf_filtered, paste0("True_Trim_", i), paste("True Class - Trim",
i))
  combined_map <- grid.arrange(predicted_map, true_map, ncol = 2)
  return(combined_map)
})

```



```

}))

# Mostrar os gráficos
lapply(mapas, print)

# Salvar os gráficos como imagens separadas
for (i in 1:4) {

  ggsave(paste0("C:/Users/user/Desktop/Dissertacao_AM_LEGAL/Analise_qualitativa/com
paracao/mapa_comparacao_trim_", i, ".png"), plot = mapas[[i]], width = 14, height = 7,
units = "in")
}

# Carregar pacotes necessários
library(ggplot2)
library(sf)
library(dplyr)
library(tidyr)
library(readxl)
library(gridExtra)

# Verificar se há algum dispositivo gráfico aberto e fechá-lo
if (!is.null(dev.list())) dev.off()

# Ler as coordenadas e os resultados
coord <-
readxl::read_excel('C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/coordenadas_de
ntro_aml.xlsx')
coord <- coord[, c(1, 2)]
coord <- cbind(coord, resultados_df)
colnames(coord)[colnames(coord) == "resultados_df$R2"] <- "R2"

# Carregar o shapefile
shapefile_path <-
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/L
imites_Amazonia_Legal_2022.shp"
sertao_sf <- st_read(shapefile_path)

# Definir e transformar o CRS
sertao_sf <- st_set_crs(sertao_sf, "EPSG:4326")
sertao_sf <- st_transform(sertao_sf, crs = "EPSG:4326")

# Separar os valores de cada trimestre em colunas separadas
coord <- coord %>%
  separate(Predicted_Class, into = paste0("Predicted_Trim_", 1:4), sep = ";") %>%
  separate(True_Class, into = paste0("True_Trim_", 1:4), sep = ";")

# Converter as colunas de classes em fatores com níveis definidos
class_levels <- c("Seco", "Chuvoso", "Extremamente Chuvoso", "Extremamente Seco")

```

```

for (i in 1:4) {
  coord[[paste0("Predicted_Trim_", i)]] <- factor(coord[[paste0("Predicted_Trim_", i)]],
levels = class_levels)
  coord[[paste0("True_Trim_", i)]] <- factor(coord[[paste0("True_Trim_", i)]], levels =
class_levels)
}

# Converter o dataframe em um objeto sf e definir o CRS
df_sf <- st_as_sf(coord, coords = c("Longitude", "Latitude"), crs = "EPSG:4326")
df_sf <- st_transform(df_sf, crs = "EPSG:4326")

# Filtrar os dados para excluir os pixels com valor zero
df_sf_filtered <- df_sf[df_sf$R2 != 0, ]

# Criar a grade de resolução 0.5 grau
grid <- st_make_grid(sertao_sf, cellsize = c(0.5, 0.5), what = "polygons")

# Associar os dados com a grade
df_sf_grid <- st_join(st_as_sf(grid), df_sf_filtered, join = st_intersects)

# Definir cores para as classes
class_colors <- c("Seco" = "blue", "Chuvoso" = "green", "Extremamente Chuvoso" =
"purple", "Extremamente Seco" = "red")

# Ajustar a posição e o tamanho do texto da legenda
legend_theme <- theme(
  legend.position = "bottom",
  legend.title = element_text(size = 10),
  legend.text = element_text(size = 8)
)

# Função para criar mapas individuais para cada trimestre
plot_map <- function(data, column, title) {
  ggplot() +
    geom_sf(data = sertao_sf, fill = "white", color = "black") +
    geom_sf(data = data, aes_string(fill = column)) +
    scale_fill_manual(values = class_colors, drop = FALSE) +
    theme_minimal() +
    labs(title = title, fill = "Classificação") +
    legend_theme
}

# Criar os mapas para cada trimestre
mapas <- lapply(1:4, function(i) {
  predicted_map <- plot_map(df_sf_grid, paste0("Predicted_Trim_", i), paste("Predicted
Class - Trim", i))
  true_map <- plot_map(df_sf_grid, paste0("True_Trim_", i), paste("True Class - Trim", i))
  combined_map <- grid.arrange(predicted_map, true_map, ncol = 2)
  return(combined_map)
})

```

```

}))

# Mostrar os gráficos
lapply(mapas, print)

# Salvar os gráficos como imagens separadas
for (i in 1:4) {

  ggsave(paste0("C:/Users/user/Desktop/Dissertacao_AM_LEGAL/Analise_qualitativa/com
paracao/mapa_comparacao_TRIM_", i, ".png"), plot = mapas[[i]], width = 14, height = 7,
units = "in")
}

```

```

library(ggplot2)
library(sf)
library(dplyr)
library(tidyr)

# Carregar o shapefile da Amazônia Legal
shapefile_path <-
"C:/Users/user/Desktop/Dissertacao_AM_LEGAL/ETL/Limites_Amazonia_Legal_2022/L
imites_Amazonia_Legal_2022.shp"
sertao_sf <- st_read(shapefile_path)

# Definir o CRS para o shapefile
sertao_sf <- st_set_crs(sertao_sf, "EPSG:4326")

# Criar a grade de resolução 0.5 grau
grid <- st_make_grid(sertao_sf, cellsize = c(0.5, 0.5), what = "polygons")

# Atribuir coordenadas ao df
df <- df %>%
  mutate(
    Predicted_Class_Trim1 = factor(Predicted_Class_Trim1, levels = c("Extremamente
Seco", "Muito Seco", "Seco", "Normalidade", "Chuvoso", "Muito Chuvoso",
"Extremamente Chuvoso")),
    Predicted_Class_Trim2 = factor(Predicted_Class_Trim2, levels = c("Extremamente
Seco", "Muito Seco", "Seco", "Normalidade", "Chuvoso", "Muito Chuvoso",
"Extremamente Chuvoso")),

```

```

Predicted_Class_Trim3 = factor(Predicted_Class_Trim3, levels = c("Extremamente
Seco", "Muito Seco", "Seco", "Normalidade", "Chuvoso", "Muito Chuvoso",
"Extremamente Chuvoso")),
Predicted_Class_Trim4 = factor(Predicted_Class_Trim4, levels = c("Extremamente
Seco", "Muito Seco", "Seco", "Normalidade", "Chuvoso", "Muito Chuvoso",
"Extremamente Chuvoso")),
True_Class_Trim1 = factor(True_Class_Trim1, levels = c("Extremamente Seco",
"Muito Seco", "Seco", "Normalidade", "Chuvoso", "Muito Chuvoso", "Extremamente
Chuvoso")),
True_Class_Trim2 = factor(True_Class_Trim2, levels = c("Extremamente Seco",
"Muito Seco", "Seco", "Normalidade", "Chuvoso", "Muito Chuvoso", "Extremamente
Chuvoso")),
True_Class_Trim3 = factor(True_Class_Trim3, levels = c("Extremamente Seco",
"Muito Seco", "Seco", "Normalidade", "Chuvoso", "Muito Chuvoso", "Extremamente
Chuvoso")),
True_Class_Trim4 = factor(True_Class_Trim4, levels = c("Extremamente Seco",
"Muito Seco", "Seco", "Normalidade", "Chuvoso", "Muito Chuvoso", "Extremamente
Chuvoso"))
)

# Criar um dataframe sf a partir da grade e associar as classes previstas e verdadeiras
grid_sf <- st_as_sf(grid)
df_sf <- st_as_sf(df, coords = c("Longitude", "Latitude"), crs = st_crs(sertao_sf))

# Associar dados da grade ao df
grid_sf <- st_join(grid_sf, df_sf, join = st_intersects)

# Função para criar um mapa
create_map <- function(data, fill_col, title) {
  ggplot() +
    geom_sf(data = sertao_sf, fill = "gray80") +
    geom_sf(data = data, aes_string(fill = fill_col), color = NA) +
    scale_fill_manual(values = c("Extremamente Seco" = "red", "Muito Seco" = "orange",
"Seco" = "yellow", "Normalidade" = "green", "Chuvoso" = "blue", "Muito Chuvoso" =
"purple", "Extremamente Chuvoso" = "pink")) +
    labs(title = title, fill = "Classificação") +
    theme_minimal()
}

# Criar os mapas
plots <- list()
for (trim in 1:4) {
  pred_col <- paste0("Predicted_Class_Trim", trim)
  true_col <- paste0("True_Class_Trim", trim)

  predicted_map <- create_map(grid_sf, pred_col, paste("Predicted Class - Trim", trim))
  true_map <- create_map(grid_sf, true_col, paste("True Class - Trim", trim))

  combined_map <- gridExtra::grid.arrange(predicted_map, true_map, ncol = 2)
}

```

```

ggsave(filename =
paste0("C:/Users/user/Desktop/Dissertacao_AM_LEGAL/Analise_qualitativa/comparacao
/mapa_comparacao_trim_", trim, ".png"), plot = combined_map, width = 16, height = 8)

```

```

plots[[trim]] <- combined_map
}

```

plots

```

library(magick)

```

```

# Função para combinar e salvar imagens

```

```

combine_and_save_images <- function(trim_num, pred_pattern, true_pattern, output_dir)
{

```

```

  # Listar arquivos no diretório

```

```

  all_files <- list.files(path = output_dir, pattern = "png$", full.names = TRUE)

```

```

  # Encontrar os arquivos que correspondem aos padrões de previsão e verdade

```

```

  pred_image_path <- grep(pred_pattern, all_files, value = TRUE)

```

```

  true_image_path <- grep(true_pattern, all_files, value = TRUE)

```

```

  if (length(pred_image_path) == 1 && length(true_image_path) == 1) {

```

```

    # Carregar as imagens

```

```

    pred_image <- image_read(pred_image_path)

```

```

    true_image <- image_read(true_image_path)

```

```

    # Combinar as imagens lado a lado

```

```

    combined_image <- image_append(c(pred_image, true_image))

```

```

    # Salvar a imagem combinada

```

```

    combined_image_path <- file.path(output_dir,

```

```

    paste0("mapa_comparacao_combined_trim_", trim_num, ".png"))

```

```

    image_write(combined_image, path = combined_image_path)

```

```

  } else {

```

```

    message("Imagens não encontradas ou múltiplas correspondências para trim ",
trim_num)

```

```

  }

```

```

}

```

```

# Diretório onde as imagens estão localizadas

```

```

output_dir <- "C:/Users/user/Desktop/Dissertacao_AM_LEGAL/Analise_qualitativa/"

```

```
# Combinar e salvar imagens para cada trimestre
combine_and_save_images(1, "mapa_predicted_trim_1.png", "mapa_true_trim_1.png",
output_dir)
combine_and_save_images(2, "mapa_predicted_trim_2.png", "mapa_true_trim_2.png",
output_dir)
combine_and_save_images(3, "mapa_predicted_trim_3.png", "mapa_true_trim_3.png",
output_dir)
combine_and_save_images(4, "mapa_predicted_trim_4.png", "mapa_true_trim_4.png",
output_dir)
```

APÊNDICE B

WORM PLOTS DE TODOS OS PIXELS DA AMAZÔNIA LEGAL BRASILEIRA

