



YASMIN BEATRIZ PEREIRA SANTANA

**IMPLEMENTAÇÃO DO MODELO $AR(1)$ X $AR(1)$ NA
ANÁLISE DE EXPERIMENTOS DE CAMPO COM
DEPENDÊNCIA ESPACIAL PARA O PACOTE
SPANOVA**

LAVRAS – MG

2026

YASMIN BEATRIZ PEREIRA SANTANA

**IMPLEMENTAÇÃO DO MODELO $AR(1)$ X $AR(1)$ NA ANÁLISE DE
EXPERIMENTOS DE CAMPO COM DEPENDÊNCIA ESPACIAL PARA
O PACOTE SPANOVA**

Dissertação apresentada à Universidade Federal de Lavras como parte das exigências do Programa de Pós-Graduação em Estatística e Experimentação Agropecuária, para obtenção do título de Mestre.

Prof. Dr. Renato Ribeiro de Lima
Orientador

**LAVRAS – MG
2026**

**Ficha Catalográfica elaborada pelo Sistema de Geração
de Ficha Catalográfica da Biblioteca Universitária da UFLA, com
dados informados pelo(a) próprio(a) autor(a).**

Santana, Yasmin Beatriz Pereira.

Implementação do modelo AR(1) X AR(1) na análise de experimentos de campo com dependência espacial para o pacote spanova / Yasmin Beatriz Pereira Santana. - 2026.

89 p. : il.

Orientador: Renato Ribeiro de Lima

Dissertação (Mestrado Acadêmico) - Universidade Federal de Lavras, 2026.
Bibliografia.

1. análise de variância. 2. erros não independentes. 3. experimentos de campo. 4. reml. I. Lima, Renato Ribeiro de. II. Universidade Federal de Lavras. III. Título.

YASMIN BEATRIZ PEREIRA SANTANA

**IMPLEMENTAÇÃO DO MODELO AR(1) X AR(1) NA ANÁLISE DE
EXPERIMENTOS DE CAMPO COM DEPENDÊNCIA ESPACIAL PARA
O PACOTE SPANOVA**

**IMPLEMENTATION OF THE AR(1) $\times$ AR(1) MODEL IN THE
ANALYSIS OF FIELD EXPERIMENTS WITH SPATIAL DEPENDENCE
FOR THE SPANOVA PACKAGE**

Dissertação apresentada à Universidade Federal de Lavras como parte das exigências do Programa de Pós-Graduação em Estatística e Experimentação Agropecuária, para obtenção do título de Mestre.

APROVADA em 02 de fevereiro de 2026.

Prof. Dr. Renato Ribeiro de Lima UFLA

Prof. Dr. Luiz Otávio de Oliveira Pala UFLA

Prof. Dr. Diogo Francisco Rossoni UEM

Prof. Dr. Renato Ribeiro de Lima
Orientador

**LAVRAS – MG
2026**

Dedico este trabalho à minha mãe, que sempre esteve comigo.

AGRADECIMENTOS

A Deus, pela vida e pela força para concluir este ciclo.

Dedico esta conquista à memória do meu pai José, cuja ausência física é preenchida diariamente pelo exemplo de força que ele deixou. Onde quer que esteja, espero que sinta orgulho desta conquista, que também é fruto das sementes que o senhor plantou..

À minha mãe, Cleide, por ser o alicerce inabalável da minha vida. Agradeço por ter caminhado ao meu lado e por ter sido o suporte fundamental em todos os momentos desta jornada. Este trabalho é o resultado de uma trajetória que só foi possível porque tive você como apoio e incentivo.

Aos meus irmãos, Pryscilla, Charles e Sofia pelo caminho compartilhado e por serem meu apoio mais sincero. Com vocês, aprendi o valor da união e encontrei o encorajamento necessário para enfrentar cada desafio desta jornada.

Aos meus avós e tios, que foram meu suporte emocional durante todo o processo. Agradeço por serem o porto seguro onde pude descansar e por acreditarem em mim nos momentos de maior pressão.

Ao meu orientador, Prof. Dr. Renato Ribeiro de Lima, pela orientação segura e pela confiança depositada. Agradeço por compartilhar seu conhecimento com tanta generosidade e por guiar este trabalho com paciência e excelência.

À Universidade Federal de Lavras, especialmente ao Departamento de Ciência Exatas, pela oportunidade, pela oportunidade de formação e por disponibilizar a infraestrutura necessária para a realização deste trabalho.

Por fim, agradeço ao Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq) e à Coordenação de Aperfeiçoamento de Pessoal de Nível Superior (CAPES). O apoio institucional e as bolsas de estudo concedidas foram os recursos indispensáveis que permitiram a dedicação integral e a viabilização desta pesquisa científica.

RESUMO

Em experimentos de campo, a análise estatística é frequentemente desafiada pela heterogeneidade espacial, onde a dependência entre as unidades experimentais representa um aspecto crítico que, se não adequadamente modelada, pode comprometer a confiabilidade das inferências. Métodos tradicionais mostram-se limitados diante da complexidade dos dados agrícolas que apresentam uma estrutura de dependência espacial. Nesse contexto, nesta dissertação apresenta-se o desenvolvimento e a integração computacional da estrutura de covariância espacial $AR(1) \times AR(1)$ ao pacote spANOVA do software R. O modelo $AR(1) \times AR(1)$ será empregado por sua capacidade de modelar eficientemente a autocorrelação residual bidimensional em arranjos regulares, visando uma representação mais realista da estrutura dos dados. O objetivo central é expandir as funcionalidades do spANOVA, permitindo que pesquisadores especifiquem este modelo em suas análises para modelos mistos. A metodologia delineada inclui o desenvolvimento computacional e a validação da função $AR(1) \times AR(1)$, a avaliação de desempenho por meio de simulações e aplicações a dados reais, utilizando a Máxima Verossimilhança Restrita (REML) para estimação de parâmetros.

Palavras-chave: análise de variância; erros não independentes; experimentos de campo; reml

ABSTRACT

In field experiments, statistical analysis is often challenged by spatial heterogeneity, in which dependence among experimental units represents a critical aspect that, if not properly modeled, may compromise the reliability of inferences. Traditional methods are limited when faced with the complexity of agricultural data that exhibit spatial dependence structures. In this context, this dissertation presents the development and computational integration of the $\text{AR}(1) \times \text{AR}(1)$ spatial covariance structure into the spANOVA package of the R software. The $\text{AR}(1) \times \text{AR}(1)$ model is employed due to its ability to efficiently model bidimensional residual autocorrelation in regular layouts, aiming at a more realistic representation of the data structure. The main objective is to expand the functionalities of spANOVA, enabling researchers to specify this model in mixed-model analyses. The proposed methodology includes the computational development and validation of the $\text{AR}(1) \times \text{AR}(1)$ function, its integration into the package, and performance evaluation through simulations and applications to real datasets, using Restricted Maximum Likelihood (REML) for parameter estimation, with the aim of improving the analysis of agricultural data with spatial dependence.

Keywords: analysis of variance; correlated errors; field experiments; restricted maximum likelihood

INDICADORES DE IMPACTO

Este trabalho apresenta impactos metodológicos e tecnológicos voltados à área acadêmica e científica, com especial contribuição para a estatística experimental e ciências agrárias. A implementação do modelo $AR(1) \times AR(1)$ no pacote spANOVA amplia o acesso a essa técnica de análise espacial por parte de pesquisadores e profissionais da experimentação agrícola, promovendo maior precisão na interpretação de dados com dependência espacial. Embora os impactos sociais e econômicos ainda sejam em potencial, a ferramenta poderá ser aplicada em estudos agropecuários, contribuindo para o uso mais eficiente de recursos experimentais. O público beneficiado inclui, estudantes, docentes e pesquisadores de Estatísticas, Agronomia, Engenharia Florestal e demais áreas de Ciências Agrárias, sendo o trabalho desenvolvido em parceria entre discente e docente no âmbito do Programa de Pós-Graduação em Estatística e Experimentação Agropecuária da UFLA. Alinha-se às áreas temáticas de: Educação (4), pela contribuição à formação acadêmica ao facilitar o uso de métodos estatísticos avançados; e de Tecnologia e Produção (7), pelo desenvolvimento de uma solução computacional inovadora voltada à análise de dados experimentais. O caráter extensionista está no potencial de disseminação da ferramenta a instituições acadêmicas e de pesquisa externas à UFLA. Adicionalmente, o estudo dialoga com o Objetivo de Desenvolvimento Sustentável 9 (Indústria, Inovação e Infraestrutura) da Agenda 2030 da ONU, por meio do desenvolvimento de uma ferramenta estatística que aprimora a análise de experimentos com dependência espacial, sendo que essa solução computacional promove o avanço da inovação científica e tecnológica.

IMPACT INDICATORS

This work presents methodological and technological impacts directed toward the academic and scientific community, with particular contribution to experimental statistics and agricultural sciences. The implementation of the $AR(1) \times AR(1)$ model in the spANOVA package broadens access to this spatial analysis technique for researchers and professionals in agricultural experimentation, promoting greater accuracy in the interpretation of data with spatial dependence. Although the social and economic impacts are still potential, the tool can be applied in agricultural and livestock studies, contributing to a more efficient use of experimental resources. The beneficiaries include students, faculty members, and researchers in Statistics, Agronomy, Forest Engineering, and other areas of Agricultural Sciences. The work was developed through a partnership between a graduate student and a faculty advisor within the Graduate Program in Statistics and Agricultural Experimentation at UFLA. The study aligns with the thematic areas of Education (4), by contributing to academic training through facilitating the use of advanced statistical methods; and Technology and Production (7), through the development of an innovative computational solution aimed at experimental data analysis. The extension-oriented character lies in the potential dissemination of the tool to academic and research institutions external to UFLA. Additionally, the study is aligned with the United Nations Sustainable Development Goal 9 (Industry, Innovation and Infrastructure) of the 2030 Agenda, through the development of a statistical tool that enhances the analysis of experiments with spatial dependence, thereby promoting advances in scientific and technological innovation

LISTA DE TABELAS

Tabela 4.1 – Testes de hipóteses globais para o efeito de tratamento sob o modelo	
$\text{AR}(1) \times \text{AR}(1)$	44
Tabela 4.2 – Parâmetros de covariância estimados via modelo autorregressivo espacial	44
Tabela 4.3 – Comparações pareadas de Wald para os contrastes entre tratamentos .	45
Tabela 4.4 – Tabela de Análise de Variância (ANOVA) do DIC	46
Tabela 4.5 – Comparação dos Critérios de Informação	46
Tabela 4.6 – Testes de hipóteses globais para o efeito de tratamento sob o modelo	
$\text{AR}(1) \times \text{AR}(1)$	48
Tabela 4.7 – Parâmetros de covariância estimados via modelo autorregressivo espacial	48
Tabela 4.8 – Médias estimadas ajustadas pelo modelo Double $\text{AR}(1)$ e agrupamento	
pelo teste de Wald ($\alpha = 0,05$)	49
Tabela 4.9 – Tabela de Análise de Variância (ANOVA) do DBC	50
Tabela 4.10 – Comparação dos Critérios de Informação	52
Tabela 4.11 – Comparação dos modelos ajustados ao conjunto de dados <code>eden.potato</code>	54

SUMÁRIO

1	INTRODUÇÃO	12
2	REFERENCIAL TEÓRICO	15
2.1	Princípios básicos da experimentação	15
2.2	Análise de variância	15
2.3	Modelagem estatística com dependência espacial	16
2.3.1	Detecção da dependência espacial	17
2.4	Modelos lineares mistos	17
2.5	Modelo $AR(1) \times AR(1)$	18
2.5.1	Formulação matemática do modelo $AR(1) \times AR(1)$	20
2.6	Modelos $AR(1) \times AR(1)$ com efeito pepita	22
2.6.1	Formulação matemática com efeito pepita	23
2.7	Estimação por máxima verossimilhança restrita - REML	23
2.8	Estimação dos efeitos fixos	25
2.9	Análise de variância utilizando o modelo $AR(1) \times AR(1)$	26
2.10	Teste de Wald	26
2.10.1	Teste Wald multivariado	27
2.10.2	Teste Wald multivariado para efeito global de fator	28
2.11	Teste da Razão de Verossimilhança (LRT)	28
2.12	Comparação pareadas de médias	29
2.13	CrITÉRIOS de informação para seleção de modelos	30
2.13.1	CrITÉRIO de informação de Akaike	30
2.13.2	CrITÉRIO de informação bayesiano	31
2.14	Principais áreas de aplicação do modelo $AR(1) \times AR(1)$	31
2.15	O pacote spANOVA	32
3	METODOLOGIA	34
3.1	Desenvolvimento da função doubleAR1	34
3.1.1	Estrutura de entrada e pré-processamento	34
3.1.2	Núcleo computacional e otimização	34
3.1.3	Módulo de inferência e robustez estatística	35
3.2	Exemplos de aplicação	35
3.2.1	Validação via simulação em DIC	35

3.2.2	Configuração experimental	36
3.2.3	Protocolo de simulação dos dados	36
3.2.4	Modelagem da superfície de tendência determinística	37
3.2.5	Objetivos da inferência e comparação de modelos	37
3.2.6	Análise para delineamento em blocos casualizados (DBC) . . .	38
3.2.6.1	Modelagem autorregressiva espacial	38
3.2.6.2	Inferência e ajuste de médias	39
3.2.7	Critérios de comparação entre modelos	39
3.3	Avaliação comparativa de desempenho	40
4	RESULTADOS E DISCUSSÃO	41
4.1	Sintaxe e detalhes da função <code>doubleAR1</code>	41
4.2	Resultados do ajuste do modelo $AR(1) \times AR(1)$ em delineamen- tos experimentais	43
4.2.1	Análise para DIC	43
4.2.1.1	Análise de significância global e verossimilhança	43
4.2.1.2	Estrutura de dependência espacial e covariância	44
4.2.1.3	Comparações múltiplas e contraste entre médias	45
4.2.1.4	Comparação com a ANOVA	46
4.2.2	Análise para DBC	47
4.2.2.1	Análise de significância global e verossimilhança	47
4.2.2.2	Estrutura de dependência espacial e covariância	48
4.2.2.3	Comparação de médias ajustadas	49
4.2.2.4	Comparação da função com a ANOVA (DBC)	50
4.2.2.5	Análise do modelo tradicional	50
4.2.2.6	Diagnóstico visual e independência dos erros	51
4.2.2.7	Comparação de ajuste e eficiência	52
4.3	Avaliação comparativa de desempenho	53
5	CONCLUSÕES	56
	REFERÊNCIAS	57
	APÊNDICES	63

1 INTRODUÇÃO

Ensaaios experimentais conduzidos em campo frequentemente apresentam heterogeneidade espacial do ambiente, o que pode comprometer a validade das análises estatísticas quando a dependência entre unidades vizinhas é ignorada (Hawinkel *et al.*, 2022). Nesse contexto, modelos que incorporam a estrutura espacial das parcelas permitem uma análise mais precisa, sendo o modelo $AR(1) \times AR(1)$, amplamente utilizado em diferentes áreas da experimentação.

Tradicionalmente, o planejamento de experimentos tem sido fundamentado nos três princípios clássicos propostos por Fisher (1935): repetição, aleatorização e controle local. A repetição é necessária na estimação da variância residual; a aleatorização busca garantir a independência dos erros; e o controle local procura reduzir os efeitos de fatores externos. Contudo, em experimentos de campo, muitas vezes esses princípios não são suficientes para garantir a independência espacial dos dados.

Assim, diversas abordagens espaciais têm sido incorporadas à estatística aplicada com o objetivo de estimar e modelar a dependência entre erros experimentais. Um dos modelos espaciais utilizados é o modelo $AR(1) \times AR(1)$, que considera a correlação residual bidimensional ao longo de duas direções: linha e coluna. Essa abordagem foi inicialmente formalizada por Grondona *et al.* (1996) e amplamente difundida por Gilmour, Cullis & Verbyla (1997), tornando-se uma ferramenta fundamental para experimentos dispostos em malha regular, uma vez que proporciona uma representação mais realista da estrutura dos dados.

A literatura demonstra a ampla aplicação e os benefícios desse modelo em diversas áreas. Em experimentos agropecuários, Piepho, Crossa & Cornelius (2015) mostram melhorias na precisão das estimativas. No melhoramento genético, trabalhos como o de Tadese & Pho (2023) destacam ganhos na acurácia na seleção de genótipos. Já em pesquisas ambientais, Chen (2013) mostra seu uso na modelagem da umidade do solo. Esses estudos são alguns exemplos de como o modelo $AR(1) \times AR(1)$ é estudado e aplicado.

Apesar da importância e da extensa utilização desse modelo, muitos pacotes estatísticos na linguagem R (R Core Team, 2026) ainda não o oferecem de forma simples e direta. Um exemplo é o pacote spANOVA (Castro *et al.*, 2024), que, embora permita análises de variância com outras estruturas espaciais, não possui o modelo $AR(1) \times AR(1)$.

em sua estrutura atual. A ausência dessa funcionalidade limita a aplicação do pacote em experimentos onde a dependência bidimensional é relevante, justificando, assim, o desenvolvimento proposto nesta dissertação.

Assim, Ensaios experimentais conduzidos em campo frequentemente apresentam heterogeneidade espacial do ambiente, o que pode comprometer a validade das análises estatísticas quando a dependência entre unidades vizinhas é ignorada (Hawinkel *et al.*, 2022). Diante desse cenário, o uso de modelos que incorporam a estrutura espacial de parcelas experimentais torna-se importante, com grande destaque para o modelo $AR(1) \times AR(1)$, amplamente utilizado em diferentes áreas da experimentação.

Tradicionalmente, o planejamento de experimentos tem sido fundamentado nos três princípios clássicos propostos por Fisher (1935): repetição, aleatorização e controle local. A repetição é necessária na estimação da variância residual; a aleatorização busca garantir a independência dos erros; e o controle local procura reduzir os efeitos de fatores externos. Contudo, em experimentos de campo, muitas vezes esses princípios não são suficientes para garantir a independência espacial desses dados.

Assim, diversas abordagens espaciais têm sido incorporadas à estatística aplicada com o objetivo de estimar e modelar a dependência entre erros experimentais. Um dos modelos espaciais utilizados é o modelo $AR(1) \times AR(1)$, que considera a correlação residual bidimensional ao longo de duas direções: linha e coluna. Essa abordagem foi inicialmente formalizada por Grondona *et al.* (1996) e amplamente difundida por Gilmour, Cullis & Verbyla (1997), tornando-se uma ferramenta fundamental para experimentos dispostos em malha regular, uma vez que proporciona uma representação mais realista da estrutura dos dados.

A literatura demonstra a ampla aplicação e os benefícios desse modelo em diversas áreas. Em experimentos agropecuários, Piepho, Crossa & Cornelius (2015) mostram melhorias na precisão das estimativas. No melhoramento genético, trabalhos como o de Tadese & Pho (2023) destacam ganhos na acurácia na seleção de genótipos. Já em pesquisas ambientais, Chen (2013) mostra seu uso na modelagem da umidade do solo. Esses estudos são alguns exemplos de como o modelo $AR(1) \times AR(1)$ é estudado e aplicado.

Apesar da importância e da extensa utilização desse modelo, muitos pacotes estatísticos na linguagem R (R Core Team, 2026) ainda não o oferecem de forma simples e direta. Um exemplo é o pacote spANOVA (Castro *et al.*, 2024), que, embora permita aná-

lises de variância com outras estruturas espaciais, não possui o modelo $AR(1) \times AR(1)$ em sua estrutura atual. A ausência dessa funcionalidade limita a aplicação do pacote em experimentos onde a dependência bidimensional é relevante, justificando, assim, o desenvolvimento proposto nesta dissertação.

Nesse contexto, o objetivo desta dissertação é implementar a estrutura de covariância bidimensional $AR(1) \times AR(1)$ no pacote **spANOVA** (Castro *et al.*, 2024), permitindo que usuários analisem experimentos em grade com dependência espacial via modelos mistos. Para tanto, o trabalho propõe o desenvolvimento e a integração dessa função como um novo módulo, assegurando a compatibilidade com os procedimentos de estimação via REML e testes F já existentes. Por fim, a implementação será validada computacionalmente por meio de simulações e aplicações em dados reais, garantindo a robustez e a eficiência da ferramenta proposta.

2 REFERENCIAL TEÓRICO

Neste capítulo, serão apresentados os principais conceitos e definições que fundamentam essa dissertação.

2.1 Princípios básicos da experimentação

Os princípios básicos da experimentação, conforme propostos por Fisher (1935), são fundamentais para garantir a validade dos resultados experimentais. Esses princípios são repetição, casualização e controle local.

A repetição consiste na aplicação de um mesmo tratamento a duas ou mais unidades experimentais. Esse princípio é fundamental para permitir a estimativa da variância do erro experimental e para aumentar a precisão das estimativas das médias. Dessa forma, obtém-se uma representação mais fiel da população-alvo, conferindo maior validade e confiabilidade aos testes de hipóteses realizados (Silva, 2024).

A casualização diz respeito aos tratamentos serem distribuídos aleatoriamente entre as unidades experimentais, garantindo que os erros sejam independentes e eliminando, assim, possíveis vieses na comparação. Esse princípio é fundamental para assegurar que os efeitos observados sejam realmente causados pelos tratamentos, evitando interferências de fatores externos (Freitas, 2022; Banzatto; Kronka, 2006).

O controle local desempenha um papel na redução da variação do erro experimental e no aumento da precisão das estimativas estatísticas. Sua aplicação consiste na criação de blocos heterogêneos entre si, mas homogêneos internamente, garantindo que a variabilidade intrablocos seja minimizada e que as estimativas dos efeitos dos tratamentos sejam mais confiáveis (Cressie, 1993).

2.2 Análise de variância

A análise de variância, ou ANOVA foi proposta por Fisher (1935). Seu principal objetivo é avaliar as possíveis fontes de variação, verificando se existem diferenças significativas entre os níveis dos fatores controlados e aleatórios.

De acordo com Banzatto & Kronka (2006), ela consiste na decomposição da variância total (e dos graus de liberdade) em partes atribuídas a causas conhecidas e independentes, fatores que são controlados, e a uma porção residual de natureza aleatória, associada a fatores não controlados.

Para a correta aplicação da ANOVA e validade do teste F, é necessário que algumas pressuposições sejam atendidas. Conforme destacado por Cochran (1947), essas suposições são:

- i) independência – os erros devem ser independentes entre si;
- ii) normalidade – os erros devem seguir uma distribuição normal com média zero e variância comum;
- iii) homogeneidade – a variância dos erros deve ser constante em todos os níveis dos tratamentos.

Essas condições garantem que as estimativas das variâncias sejam válidas e que o teste F possa ser utilizado para verificar a existência de diferenças significativas entre as médias dos tratamentos.

2.3 Modelagem estatística com dependência espacial

Em experimentos de campo, especialmente em áreas agrícolas, é comum que unidades experimentais próximas apresentem respostas semelhantes devido à influência de fatores ambientais e características do solo. Essa autocorrelação espacial viola a suposição de independência dos erros, fundamental para modelos estatísticos clássicos (Legendre, 1993; Otto; Fassò; Maranzano, 2024).

A presença da dependência espacial distorce os valores-p obtidos na análise de variância, subestimando a variância residual e aumentando a ocorrência de erros do Tipo I, onde efeitos são considerados significativos mesmo sem evidência real (Darghan *et al.*, 2024). Dessa forma, a simples aleatorização das unidades experimentais não é suficiente para mitigar completamente esse problema, o que pode comprometer a validade das conclusões (Resende; Thompson, 2004).

Nesse contexto, abordagens que incorporam explicitamente a estrutura espacial dos dados mostram-se mais apropriadas para descrever com precisão os dados observados (Gaetan; Guyon, 2010; Dumelle; Higham; Hoef, 2023).

2.3.1 Detecção da dependência espacial

Uma vez reconhecida a importância da dependência espacial, torna-se essencial detectar sua presença nos dados experimentais antes do ajuste de modelos. De acordo com Diggle & Ribeiro (2007), é possível detectar indícios de tendência espacial em uma análise exploratória dos dados, por meio de gráficos de quartis e gráficos da variável resposta em relação às suas coordenadas. Essas ferramentas permitem observar padrões sistemáticos que violam a suposição de independência dos erros.

Antes da escolha de um modelo com estrutura espacial, é essencial investigar se há presença de dependência entre unidades experimentais adjacentes. Como destaca Fortin & Legendre (2020), a autocorrelação espacial residual pode ser diagnosticada por meio da análise gráfica dos resíduos de modelos clássicos, como a ANOVA, indicando agrupamentos ou gradientes espaciais.

Além da inspeção visual, medidas estatísticas formais também são utilizadas, como o Índice de Moran, amplamente empregado para quantificar a autocorrelação espacial ao avaliar a semelhança entre valores observados em locais próximos (Chen, 2013; Dumelle; Higham; Hoef, 2023).

Outra ferramenta que pode ser utilizada é o variograma experimental, que segundo Scalón (2024) descreve a variabilidade dos dados em função da distância entre pares de observações, permitindo visualizar o grau e o alcance da dependência espacial. Esta técnica é empregada na Geoestatística como diagnóstico preliminar da estrutura espacial dos dados.

A combinação de métodos gráficos, estatísticos e geoestatísticos fornece uma base sólida para a detecção de dependência espacial, orientando a necessidade de incorporar estruturas de correlação nos modelos estatísticos subsequentes.

2.4 Modelos lineares mistos

Os modelos lineares mistos representam uma extensão dos modelos lineares clássicos, permitindo a análise conjunta de efeitos fixos e aleatórios. Em termos práticos, enquanto os efeitos fixos (β) descrevem a média dos tratamentos de interesse, os efeitos

aleatórios ($\mathbf{u}$) modelam fontes de variação provenientes do delineamento experimental, como blocos e parcelas (Montgomery, 2017; Resende; Duarte, 2007).

A estrutura matricial que define o modelo é expressa pela equação 2.1

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{u} + \boldsymbol{\epsilon}, \quad (2.1)$$

em que $\mathbf{y}$ é o vetor de dados observados; $\mathbf{X}$ e $\mathbf{Z}$ são matrizes de incidência que conectam as observações aos seus respectivos efeitos; $\boldsymbol{\beta}$ é o vetor de efeitos fixos a serem estimados; $\mathbf{u} \sim N(\mathbf{0}, \mathbf{G})$ representa o vetor de efeitos aleatórios e $\boldsymbol{\epsilon} \sim N(\mathbf{0}, \mathbf{R})$ representa o erro residual.

O diferencial dessa modelagem reside na estrutura da variância total dos dados, denotada pela matriz $\mathbf{V}$. Diferente da ANOVA clássica, onde se assume independência entre as observações, nos modelos mistos a variância é dada por $\mathbf{V} = \mathbf{ZGZ}^\top + \mathbf{R}$ (Littel *et al.*, 2006; Patterson; Thompson, 2023). Nessa composição, a matriz $\mathbf{G}$ gerencia a variabilidade dos efeitos aleatórios, enquanto a matriz $\mathbf{R}$ lida com a variabilidade dos resíduos, permitindo que o modelo considere a correlação entre as observações

Uma vez que a dependência espacial tenha sido detectada, essa formulação permite que a matriz de variância-covariância incorpore as correlações entre observações vizinhas. Essa flexibilidade é indispensável para garantir que as inferências sobre os tratamentos não sejam tendenciosas pela proximidade física das parcelas (Piepho *et al.*, 2022).

Por conta disso, os modelos lineares mistos possibilitam a especificação de estruturas de covariância que tornam a análise robusta em ambientes onde os pressupostos da ANOVA não são plenamente atendidos (Stroup, 2018). Neste contexto, um modelo que se destaca por sua eficácia na captura da dependência bidimensional em experimentos agrícolas é o $\text{AR}(1) \times \text{AR}(1)$.

2.5 Modelo $\text{AR}(1) \times \text{AR}(1)$

Modelos autorregressivos (AR) partem do princípio de que a resposta observada em um determinado local, Y_i , depende não apenas das variáveis explicativas associadas àquela localização, mas também dos valores das respostas observadas em localidades vizi-

nhas (Cressie, 1993). Dessa forma, a estrutura autorregressiva desses modelos exige uma definição explícita da vizinhança espacial envolvida.

A variabilidade espacial em experimentos pode ser abordada por diversas metodologias. Considerando o contexto da análise de séries temporais aplicadas ao espaço, Gleeson & Cullis (1987) sugeriram o uso de um modelo autorregressivo de primeira ordem, $AR(1)$, para modelar a dependência espacial dos resíduos ao longo de uma única dimensão, como linhas ou colunas de uma grade espacial.

Na sua forma unidimensional, o modelo $AR(1)$ descreve a correlação entre observações adjacentes ao longo de uma direção, partindo do pressuposto de que essa influência diminui com a distância. Segundo Grzesiek (2021), processos autorregressivos pertencem à classe de processos de memória curta, caracterizados por uma autocorrelação que decresce exponencialmente.

Para modelar a dependência espacial bidimensional, destaca-se o modelo autorregressivo separável de primeira ordem, conhecido como $AR(1) \times AR(1)$. Esse modelo foi concebido como uma extensão do modelo $AR(1)$ unidimensional por Cullis & Gleeson (1991). Posteriormente, Grondona *et al.* (1996) formalizaram a extensão bidimensional do modelo por meio da formulação separável baseada no produto de Kronecker entre as autocorrelações nas direções de linha e coluna, estabelecendo a estrutura conhecida como $AR(1) \times AR(1)$. Em seguida, Gilmour, Cullis & Verbyla (1997) consolidaram a aplicação dessa estrutura nos experimentos agrícolas, destacando sua capacidade de capturar múltiplas fontes de variação espacial e reforçando sua adoção na modelagem de dados organizados em malha regular, isto é, dispostos em linhas e colunas com espaçamentos constantes.

Desse modo, o $AR(1) \times AR(1)$ é um modelo versátil por sua capacidade de representar diferentes padrões de variação espacial observados em experimentos de campo. Segundo Hoefler *et al.* (2020), essa estrutura permite capturar três fontes principais de variação: tendências globais, variações locais e variação extrínseca ou aleatória, relacionadas a fatores não estruturais. As tendências globais referem-se a mudanças ao longo de toda a área experimental; as variações locais seriam as pequenas oscilações entre unidades vizinhas; e os componentes aleatórios são atribuídos a fatores que não estão em sua estrutura. Essa flexibilidade torna o modelo indicado para experimentos com forte estrutura

espacial, sendo frequentemente adotados em modelagem de dados agrícolas (Hawinkel *et al.*, 2022).

2.5.1 Formulação matemática do modelo $\text{AR}(1) \times \text{AR}(1)$

O modelo $\text{AR}(1) \times \text{AR}(1)$ é construído sobre a estrutura de um modelo linear geral, que pode ser expresso na forma matricial como na Equação 2.2

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}, \quad (2.2)$$

em que $\mathbf{y}$ é o vetor coluna ($n \times 1$) das observações da variável resposta, sendo n o número total de unidades experimentais; $\mathbf{X}$ é a matriz de incidência de dimensão $n \times p$, em que p representa o número de parâmetros de efeitos fixos a serem estimados; $\boldsymbol{\beta}$ é o vetor ($p \times 1$) dos coeficientes dos efeitos fixos e $\boldsymbol{\varepsilon}$ representa o vetor de erros aleatórios, com estrutura de dependência espacial modelada pela matriz de covariância.

O diferencial do modelo $\text{AR}(1) \times \text{AR}(1)$ em relação ao modelo linear simples é a suposição específica sobre a distribuição e a estrutura de covariância dos erros aleatórios ($\boldsymbol{\varepsilon}$). Para este modelo, assume-se que os erros seguem uma distribuição normal multivariada, caracterizada por um vetor de médias nulo e uma matriz de covariância particular, conforme representado pela Equação 2.3

$$\boldsymbol{\varepsilon} \sim N(\mathbf{0}, \mathbf{V}), \quad (2.3)$$

em que o valor esperado do vetor de erros é nulo, implicando a ausência de viés, e $\mathbf{V}$ representa a matriz de covariância dos erros, definida pela Equação 2.4

$$\mathbf{V} = \sigma^2(\mathbf{R}_{\text{linha}} \otimes \mathbf{R}_{\text{coluna}}). \quad (2.4)$$

Em que σ^2 representa a variância total dos erros. A matriz de covariância $\mathbf{V}$ é construída a partir do produto de Kronecker ($\otimes$) entre as matrizes de correlação autor-regressiva de primeira ordem, $\mathbf{R}_{\text{linha}}$ e $\mathbf{R}_{\text{coluna}}$. Esse produto de matrizes é usado para combinar as estruturas de dependência unidimensionais das linhas e colunas em uma única matriz de covariância bidimensional. Cada uma das matrizes de correlação segue o padrão

AR(1), em que a correlação entre duas observações separadas por uma distância d é dada por ρ^d , sendo ρ o coeficiente de autocorrelação.

A matriz de correlação associada às linhas, com dimensão $L \times L$, possui a forma:

$$R_{\text{linha}} = \begin{bmatrix} 1 & \rho_L & \rho_L^2 & \cdots & \rho_L^{L-1} \\ \rho_L & 1 & \rho_L & \cdots & \rho_L^{L-2} \\ \rho_L^2 & \rho_L & 1 & \cdots & \rho_L^{L-3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho_L^{L-1} & \rho_L^{L-2} & \rho_L^{L-3} & \cdots & 1 \end{bmatrix}$$

A matriz das colunas, com dimensão $C \times C$, assume uma estrutura análoga, substituindo-se o parâmetro de linha (ρ_L) pelo de coluna (ρ_C), e a dimensão de linha (L) pela de coluna (C). A sua forma é:

$$R_{\text{coluna}} = \begin{bmatrix} 1 & \rho_C & \rho_C^2 & \cdots & \rho_C^{C-1} \\ \rho_C & 1 & \rho_C & \cdots & \rho_C^{C-2} \\ \rho_C^2 & \rho_C & 1 & \cdots & \rho_C^{C-3} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \rho_C^{C-1} & \rho_C^{C-2} & \rho_C^{C-3} & \cdots & 1 \end{bmatrix}$$

Essa construção garante que observações vizinhas apresentem correlações mais fortes do que observações distantes, conforme esperado em dados espaciais. A matriz de covariância $\mathbf{V}$, que terá dimensão $(n \times n)$, em que $n = L \times C$ e corresponde ao número total de observações, é o resultado do produto de Kronecker de $\mathbf{R}_{\text{linha}}$ e $\mathbf{R}_{\text{coluna}}$.

A principal vantagem desse modelo bidimensional é a flexibilidade para estimar coeficientes de autocorrelação distintos para a linha (ρ_L) e a coluna (ρ_C), uma característica crucial para a modelagem da anisotropia, fenômeno que ocorre quando a dependência espacial não é uniforme em todas as direções (Cressie, 1993; Müller; Waagepetersen, 2023).

Além disso, a estrutura gerada pelo produto de Kronecker mantém a simetria e a positividade definida da matriz de covariância total, o que garante a validade estatística do modelo e viabiliza a aplicação de métodos como máxima verossimilhança restrita (REML) para a estimação dos parâmetros (Cressie, 1993).

É importante notar que, em algumas situações, pode haver uma porção da variabilidade residual que não apresenta estrutura espacial detectável, o efeito pepita. Essa extensão do modelo, incorporando esse componente, será abordada na próxima seção.

2.6 Modelos $AR(1) \times AR(1)$ com efeito pepita

Apesar da robustez do modelo $AR(1) \times AR(1)$ na modelagem de dependência espacial bidimensional, é comum que parte da variabilidade residual não seja capturada pela estrutura autorregressiva principal (Rodríguez-Álvarez *et al.*, 2016; Müller; Schuhmacher; Mateu, 2024).

Esse componente adicional de variância, conhecido como efeito pepita (τ^2), representa uma descontinuidade na origem do semivariograma, indicando a presença de variabilidade em escalas menores que a distância mínima entre amostras ou associada a erros de medição, heterogeneidade microespacial ou ruído experimental (Deutsch, 2019). O efeito pepita reflete a parte da variância que não apresenta estrutura espacial detectável, sendo modelado como um componente aleatório não correlacionado.

Em modelos $AR(1) \times AR(1)$, a presença de um efeito pepita significativo pode indicar que, além da autocorrelação espacial capturada pelas estruturas autorregressivas, há uma parcela de ruído que deve ser considerada separadamente. Ignorar esse componente pode levar à subestimação da variância residual e comprometer a acurácia das estimativas (Silva; Lima; Andrade, 2020).

Tang, Zhang & Banerjee (2021) destacam que o valor do efeito pepita pode variar conforme o tipo de variável analisada, a escala de suporte amostral, a precisão dos instrumentos utilizados e o método de amostragem. Sendo assim, a variabilidade associada a efeitos externos e à heterogeneidade em escalas muito pequenas torna-se essencial para garantir a consistência estatística dos modelos espaciais.

2.6.1 Formulação matemática com efeito pepita

Para acomodar a variabilidade adicional do efeito pepita, os erros aleatórios do modelo (ε) são decompostos em dois componentes: um erro espacialmente correlacionado, representado por ξ , e um erro aleatório independente que reflete o efeito pepita, representado por η . Dessa forma, o vetor de erros é expresso pela Equação 2.5

$$\varepsilon = \xi + \eta. \quad (2.5)$$

Em termos matriciais, a estrutura de covariância total dos erros aleatórios (ε) de um modelo $\text{AR}(1) \times \text{AR}(1)$ com efeito pepita é estendida e representada pela Equação 2.6

$$\mathbf{V} = \sigma^2(\mathbf{R}_{\text{linha}} \otimes \mathbf{R}_{\text{coluna}}) + \tau^2 \mathbf{I}, \quad (2.6)$$

em que $\mathbf{V}$ é a matriz de covariância total dos erros residuais; σ^2 denota a variância associada à estrutura espacial correlacionada; $\mathbf{R}_{\text{linha}}$ e $\mathbf{R}_{\text{coluna}}$ representam as matrizes de autocorrelação $\text{AR}(1)$ nas direções de linha e coluna, respectivamente; $\otimes$ indica o produto de Kronecker; τ^2 representa a variância do efeito pepita e $\mathbf{I}$ é a matriz identidade de ordem n , sendo n o número total de unidades experimentais.

A inclusão do termo $\tau^2 \mathbf{I}$ representa a parte da variabilidade que não apresenta dependência espacial aparente, permitindo que o modelo acomode de forma mais realista a heterogeneidade observada nos dados (Rodríguez-Álvarez *et al.*, 2016; Diggle; Ribeiro, 2007).

2.7 Estimação por máxima verossimilhança restrita - REML

A estimação por máxima verossimilhança restrita (*Restricted Maximum Likelihood* - REML) é amplamente utilizada para inferir componentes de variância em modelos mistos, pois reduz vieses nas estimativas, tornando-as mais confiáveis (Laporte; Charcosset; Mary-Huard, 2022; Resende *et al.*, 2018).

Diferentemente da máxima verossimilhança tradicional, o REML considera apenas a parte da verossimilhança que não depende dos efeitos fixos, proporcionando melhores estimativas, especialmente em amostras pequenas ou moderadas. Essa abordagem é fun-

damental para evitar a subestimação da variância residual e melhorar a robustez dos testes estatísticos (Ferreira, 2024).

No contexto dos modelos espaciais $\text{AR}(1) \times \text{AR}(1)$ a estimação por REML possibilita o ajuste consistente dos parâmetros da estrutura de covariância, tanto em modelos que assumem somente a dependência espacial quanto naqueles que incorporam a componente do ruído adicional, o efeito pepita (Chen; Genton; Sun, 2021).

Os parâmetros de covariância, objetos de estimação via REML, são agrupados no vetor $\boldsymbol{\theta}$. A composição deste vetor varia conforme a complexidade da estrutura de erro adotada: para o modelo sem efeito pepita, define-se $\boldsymbol{\theta} = (\sigma^2, \rho_r, \rho_c)^\top$; já para o modelo completo, inclui-se o componente de variação local, resultando em $\boldsymbol{\theta} = (\sigma^2, \tau^2, \rho_r, \rho_c)^\top$. Os componentes que integram essas definições são detalhados a seguir:

- σ^2 : variância da componente espacialmente correlacionada. No modelo sem efeito pepita, esta é a variância total dos erros;
- τ^2 : variância do efeito pepita;
- ρ_{linha} : coeficiente de autocorrelação espacial na direção das linhas;
- ρ_{coluna} : coeficiente de autocorrelação espacial na direção das colunas.

Os efeitos fixos do modelo ($\boldsymbol{\beta}$) são estimados em uma etapa subsequente, após a estimação desses componentes de variância.

A função de log-verossimilhança residual (REML) é expressa pela 2.7,

$$\ell_R(\boldsymbol{\theta}) = -\frac{1}{2} \left[(n-p) \log(2\pi) + \log |\mathbf{V}| + \log |\mathbf{X}^\top \mathbf{V}^{-1} \mathbf{X}| + \mathbf{y}^\top \mathbf{P} \mathbf{y} \right], \quad (2.7)$$

em que:

- n é o número total de observações;
- p é o número de efeitos fixos no modelo;
- A matriz de variância-covariância $\mathbf{V}$ é definida por: $\mathbf{V} = \sigma^2(\mathbf{R}_{\text{linha}} \otimes \mathbf{R}_{\text{coluna}})$ no modelo sem efeito pepita e por $\mathbf{V} = \sigma^2(\mathbf{R}_{\text{linha}} \otimes \mathbf{R}_{\text{coluna}}) + \tau^2 \mathbf{I}$ no modelo com efeito pepita.
- $\mathbf{P} = \mathbf{V}^{-1} - \mathbf{V}^{-1} \mathbf{X} (\mathbf{X}^\top \mathbf{V}^{-1} \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{V}^{-1}$ é a matriz de projeção dos resíduos.

A maximização da função apresentada na Equação (2.7) é feita por métodos numéricos iterativos. O algoritmo especificamente empregado é o Quasi-Newton, utilizando sua versão eficiente L-BFGS-B (Broyden-Fletcher-Goldfarb-Shanno com memória limitada e limites de caixa) (Nocedal; Wright, 2006; Bonat *et al.*, 2024).

O método Quasi-Newton é utilizado devido à sua rápida convergência em problemas de otimização não lineares. Ele aproxima iterativamente a Hessiana da log-verossimilhança restrita, permitindo estimar os componentes de variância que definem a matriz de covariância $\mathbf{V}$ (Bonat; Jørgensen, 2016).

2.8 Estimação dos efeitos fixos

Uma vez que os parâmetros da estrutura de covariância são estimados, os efeitos fixos do modelo (β) são calculados em um passo subsequente.

O cálculo dos efeitos fixos é uma aplicação do método de Mínimos Quadrados Generalizados (GLS). O GLS é uma extensão do método de Mínimos Quadrados Ordinários, sendo especialmente útil quando os erros do modelo são correlacionados ou apresentam variâncias desiguais, como ocorre em dados com estrutura espacial. Ao utilizar a matriz de covariância estimada ($\mathbf{V}$), o GLS ajusta o cálculo para ponderar adequadamente as observações, garantindo que a estimativa dos efeitos fixos seja o Melhor Estimador Linear Não Viesado (BLUE) (Aitken, 1935; Moraga, 2023).

O vetor de efeitos fixos β é então estimado pela equação 2.8:

$$\hat{\beta} = (\mathbf{X}^\top \mathbf{V}^{-1} \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{V}^{-1} \mathbf{y}. \quad (2.8)$$

Nessa fórmula:

- $\mathbf{X}$ é a matriz de incidência;
- $\mathbf{V}$ é a matriz de variância-covariância, já estimada na Etapa 1;
- $\mathbf{y}$ é o vetor das observações.

2.9 Análise de variância utilizando o modelo $AR(1) \times AR(1)$

A ANOVA tradicional não é aplicável neste modelo, uma vez que este método é baseado no pressuposto de independência dos resíduos, uma condição que é sistematicamente violada em experimentos de campo devido à dependência espacial entre as parcelas. A utilização da ANOVA nessas condições leva a conclusões estatísticas incorretas (Cullis; Gleeson, 1991; Ferreira; Dias; Oliveira, 2021).

O modelo $AR(1) \times AR(1)$ corrige a dependência espacial por meio da modelagem explícita da matriz de covariância dos erros ($\mathbf{V}$), cuja estimação é realizada por Máxima Verossimilhança Restrita (REML) (Patterson; Thompson, 1971; Maestrini; Hui; Welsh, 2025).

Com a matriz $\mathbf{V}$ corrigida e a validade estatística garantida, a inferência dos efeitos fixos é realizada por meio de testes análogos ao Teste F da ANOVA. O Teste de Wald Multivariado e o Teste da Razão de Verossimilhança (LRT) são utilizados para avaliar o efeito global do Tratamento, e o Teste de Wald Univariado é empregado para as Comparações Pareadas (*post hoc*) (Godambe, 1960; Bonat *et al.*, 2024).

2.10 Teste de Wald

O teste de Wald é amplamente utilizado para avaliar a significância dos parâmetros em modelos lineares mistos, especialmente aqueles que apresentam dependência espacial. De acordo com Ywata & Albuquerque (2011), sua formulação baseia-se nas propriedades assintóticas dos estimadores de máxima verossimilhança, tornando-o aplicável a uma ampla variedade de modelos estatísticos.

Segundo Resende *et al.* (2018), em análises envolvendo estruturas complexas de covariância, o teste de Wald aplicado a modelos ajustados via REML tem se consolidado como uma alternativa robusta ao teste F .

No caso do modelo $AR(1) \times AR(1)$, o teste de Wald é particularmente relevante, pois considera a autocorrelação espacial ao calcular os erros-padrões, tornando as inferências estatísticas mais precisas. Esse ajuste é essencial em experimentos de campo, nos quais há dependência entre parcelas vizinhas (Gilmour; Cullis; Verbyla, 1997).

A significância dos efeitos fixos é avaliada pelo teste de Wald, cuja estatística (W) baseia-se no quadrado da razão entre a estimativa do parâmetro e seu respectivo erro-padrão, conforme a Equação 2.9:

$$W = \left(\frac{\hat{\beta}}{\text{EP}(\hat{\beta})} \right)^2, \quad (2.9)$$

em que $\hat{\beta}$ representa o valor estimado do parâmetro e $\text{EP}(\hat{\beta})$ é o seu erro-padrão. Além disso, sob a hipótese nula, essa estatística segue uma distribuição qui-quadrado com 1 grau de liberdade.

2.10.1 Teste Wald multivariado

O Teste de Wald Multivariado, uma generalização do Teste de Wald, para avaliar o efeito global de fatores com múltiplos níveis, como o Tratamento, que é representado por um vetor de vários parâmetros. Este teste constitui a principal ferramenta de inferência estatística para avaliar hipóteses lineares complexas no contexto de dependência espacial (McCulloch; Searle; Neuhaus, 2008; Mendes; Bonat; Jørgensen, 2020).

Seu objetivo é possibilitar a avaliação de hipóteses sobre um subconjunto de parâmetros de interesse, denotado por θ^* . No contexto experimental, este subconjunto compreende os coeficientes diferenciais que se deseja testar (Mello *et al.*, 2021).

Para viabilizar essa análise, utiliza-se a submatriz de informação J^* , que representa uma partição da Matriz de Informação de Godambe. Essa matriz é derivada da curvatura da função de estimação e a J^* extrai a variabilidade associada aos elementos de θ^* , isolando-os das demais fontes de variação do modelo. A utilização dessa estrutura permite uma avaliação de componentes específicos do preditor linear, mesmo sob dependência espacial ou temporal (Bonat; Jørgensen, 2021).

Dessa forma, a robustez do Teste de Wald Multivariado advém diretamente do uso da Matriz de Godambe, o que assegura a validade da inferência independentemente da estrutura de correlação modelada (Saeidi; Abbaszadeh; Jazi, 2021).

O método garante que a estatística W siga assintoticamente a distribuição Qui-quadrado (χ^2), permitindo a determinação da significância global do fator em questão (Liang; Zeger, 1986; Koul; Song; Qian, 2021).

2.10.2 Teste Wald multivariado para efeito global de fator

O Teste de Wald generalizado pode ser empregado com o propósito de avaliar a significância global do fator Tratamento, especialmente quando os dados apresentam dependência espacial (Wald, 1943; Ahsan; Baksh; Hoque, 2022).

Neste caso, a análise é conduzida sob um par de hipóteses lineares:

$$H_0 : \mathbf{L}\theta^* = \mathbf{0} \quad \text{contra} \quad H_1 : \mathbf{L}\theta^* \neq \mathbf{0}. \quad (2.10)$$

Nesta formulação, a hipótese nula (H_0) postula a ausência de efeito do fator Tratamento. A Matriz de Especificação da Hipótese ($\mathbf{L}$) é utilizada para isolar e testar conjuntamente todos os s parâmetros de contraste associados aos níveis do fator, denotados por $\beta_2, \beta_3, \dots, \beta_k$. Esses coeficientes representam os contrastes de cada nível em relação ao nível de referência (β_1), que é absorvido pelo intercepto do modelo. Este isolamento garante que o teste avalie o efeito do fator de forma global, verificando se as diferenças entre os níveis são estatisticamente nulas, de forma análoga ao teste F tradicional aplicado em modelos lineares clássicos (Freitas *et al.*, 2022; Searle, 2016).

A estatística W representa a métrica central do teste, mensurando a discrepância quadrática entre as estimativas dos coeficientes e o vetor nulo ($\mathbf{0}$), ponderada pela variabilidade corrigida do modelo. O resultado do teste (P -valor) é obtido comparando-se o valor de W à distribuição Qui-quadrado (χ^2), com graus de liberdade correspondentes ao número de restrições impostas por $\mathbf{L}$. Este procedimento fornece uma conclusão estatística sobre a significância global do tratamento, sendo devidamente corrigido pela estrutura de dependência espacial via Matriz de Informação de Godambe (J^*) (Godambe, 1960; Bonat *et al.*, 2024)

2.11 Teste da Razão de Verossimilhança (LRT)

O Teste da Razão de Verossimilhança (LRT) é empregado para comparar modelos estatísticos aninhados, permitindo determinar se a inclusão de determinados parâmetros resulta em uma melhoria significativa no ajuste do modelo (Glas; Verhelst, 1995).

A estatística do LRT é definida pela diferença entre os logaritmos das verossimilhanças de dois modelos, conforme a Equação 2.11 baseada em Wilks (1938):

$$\text{LRT} = -2\log L_0 - (-2\log L_1), \quad (2.11)$$

em que $\log L_0$ representa o logaritmo da verossimilhança do modelo reduzido e $\log L_1$ corresponde ao logaritmo da verossimilhança do modelo completo.

O LRT pode ser aplicado especificamente para avaliar a significância global dos efeitos fixos do fator de tratamento (Zuur *et al.*, 2009; West; Welch; Galecki, 2022). Nesse contexto, o modelo completo (L_1) inclui o fator tratamento, enquanto o modelo reduzido (L_0) retira esse fator, mantendo apenas o intercepto e a correção espacial $AR(1) \times AR(1)$, permitindo testar se a inclusão do tratamento melhora estatisticamente o ajuste do modelo.

Para garantir a validade estatística na comparação de modelos com diferentes efeitos fixos, o ajuste deve ser realizado via Máxima Verossimilhança. Essa escolha é necessária porque a REML é inadequada para contrastar modelos cujas partes fixas diferem entre si (Pinheiro; Bates, 2000). Sob a hipótese nula, a estatística resultante segue uma distribuição qui-quadrado com graus de liberdade equivalentes à diferença no número de parâmetros entre os modelos comparados.

A confirmação da significância global do tratamento por meio do LRT atua como um pré-requisito para as etapas seguintes da análise. Esse resultado fornece o embasamento estatístico necessário para a aplicação de testes de comparação múltipla de médias e a interpretação de contrastes individuais (Sá *et al.*, 2023; Diggle; Rowlingson; Su, 2021).

2.12 Comparação pareadas de médias

Com a significância global do fator Tratamento pelo Teste de Wald Multivariado, uma análise *post hoc* é conduzida para determinar quais diferenças específicas entre os níveis do tratamento são estatisticamente significativas (Freitas, 2022).

Para este propósito, o Teste de Wald generalizado é novamente empregado, mas desta vez em sua forma univariada, executando as comparações pareadas entre todos os

pares de tratamentos. Este procedimento atua como o análogo funcional de testes de comparações múltiplas *post hoc* em modelos de regressão (Mendes; Freitas; Bonat, 2023).

A estatística do Teste de Wald para a comparação entre os tratamentos i e j , baseada na formulação originária de Wald (1943) é dada pela equação 2.12:

$$W = \frac{(\hat{\beta}_i - \hat{\beta}_j)^2}{\text{Var}(\hat{\beta}_i - \hat{\beta}_j)} \quad (2.12)$$

Em que os termos $\hat{\beta}_i$ e $\hat{\beta}_j$ representam os estimadores dos efeitos fixos associados aos tratamentos i e j , respectivamente. Esses coeficientes são obtidos via REML e refletem a resposta média estimada para cada nível do fator de tratamento após o ajuste do modelo. O numerador $(\hat{\beta}_i - \hat{\beta}_j)^2$ estabelece o quadrado da diferença entre essas estimativas, enquanto o denominador $\text{Var}(\hat{\beta}_i - \hat{\beta}_j)$ indica a variância desse contraste, calculada a partir da matriz de precisão que contempla a estrutura de dependência espacial $AR(1) \times AR(1)$.

2.13 Critérios de informação para seleção de modelos

A seleção de modelos estatísticos com dependência espacial leva em conta não apenas a significância dos parâmetros, mas também o equilíbrio entre a qualidade de ajuste e a complexidade do modelo.

Sendo assim, critérios como o Critério de Informação de Akaike (AIC) e o Critério de Informação Bayesiano (BIC) são utilizados para avaliar a pertinência da escolha de um modelo, uma vez que ambos consideram tanto o ajuste dos dados quanto a parcimônia da estrutura proposta.

2.13.1 Critério de informação de Akaike

O AIC, proposto por Akaike (1974), é utilizado para seleção de modelos estatísticos, especialmente em contextos onde diferentes estruturas de variância estão sendo comparadas. Esse critério busca equilibrar a qualidade de ajuste do modelo com sua complexidade, penalizando a inclusão de parâmetros adicionais para evitar sobreajuste.

O AIC é definido por:

$$\text{AIC} = 2k - 2\ln(\hat{L}) \quad (2.13)$$

em que k representa o número de parâmetros estimados no modelo e $\hat{L}$ é o valor da verossimilhança máxima. Quanto menor o valor de AIC, melhor o modelo é considerado, em comparação a outros modelos ajustados ao mesmo conjunto de dados.

2.13.2 Critério de informação bayesiano

O BIC, foi proposto por Schwarz (1978), sendo definido por:

$$\text{BIC} = k\ln(n) - 2\ln(\hat{L}), \quad (2.14)$$

em que n representa o número de observações, k o número de parâmetros estimados e $\hat{L}$ o valor da verossimilhança máxima. A presença do termo $\ln(n)$ torna a penalização agravada à medida que o tamanho da amostra aumenta, o que favorece a simplicidade dos modelos, principalmente em grandes amostras.

Segundo Resende (2014), o BIC tende a selecionar modelos mais parcimoniosos do que o AIC. Assim, é mais adequado quando se deseja o maior controle sobre a complexidade do modelo ajustado, ou seja, penalizando com mais intensidade modelos complexos.

2.14 Principais áreas de aplicação do modelo $\text{AR}(1) \times \text{AR}(1)$

O modelo $\text{AR}(1) \times \text{AR}(1)$ é amplamente utilizado em áreas que envolvem dados organizados em grades espaciais bidimensionais, especialmente nas áreas de experimentação agropecuária, estudos florestais, melhoramento genético de plantas e pesquisas ambientais voltadas à modelagem de recursos naturais.

Essa modelagem contribui para corrigir autocorrelações espaciais e melhorar a comparação entre tratamentos, conforme demonstrado por Selle, Steinsland & Hickey (2019), sendo implementada computacionalmente por meio do método *Integrated Nested Laplace Approximation* (INLA) (Rue; Martino; Chopin, 2009) no ambiente estatístico R (R Core Team, 2026), enquanto Piepho, Crossa & Cornelius (2015) discutiram desafios

técnicos na estimação de parâmetros espaciais em modelos AR(1), propondo estratégias para lidar com problemas de convergência e interpretação em dados agrícolas.

Em programas de melhoramento genético de plantas, o modelo tem sido eficaz para modelar a variabilidade espacial e melhorar a acurácia das estimativas genéticas. Tadese & Pho (2023) aplicaram diferentes métodos de seleção de modelos espaciais em ensaios com sorgo na Etiópia, destacando os benefícios da abordagem autorregressiva. Já Silva, Xavier & Faria (2021) avaliaram a integração de modelos genéticos e espaciais, concluindo que o uso do $AR(1) \times AR(1)$ reduziu vieses e aumentou a precisão das estimativas.

No contexto de análises agrícolas voltadas à seleção de genótipos, Bernardeli & Borém (2021) evidenciaram que a adoção do modelo resultou em maior ganho genético e mais precisão na identificação de cultivares superiores em experimentos com soja. Isso reforça o valor da modelagem espacial como ferramenta complementar ao processo de melhoramento.

No campo dos estudos florestais, o modelo $AR(1) \times AR(1)$ é utilizado para avaliar a estrutura genética e os padrões de crescimento de árvores. Um exemplo é o trabalho de Lee, Kim & Kang (2022), que utilizou a modelagem espacial combinada com a reconstrução de pedigree em testes com *Larix kaempferi*, para uma maior precisão nas estimativas de parâmetros genéticos.

Por fim, em estudos ambientais e de recursos naturais, o modelo tem sido utilizado para capturar padrões espaciais em fenômenos ecológicos, como a distribuição da umidade do solo. Ajami H e Sharma (2018) compararam diferentes abordagens para desagregação de dados espaciais, incluindo modelos autorregressivos, e comprovaram sua efetividade em contextos que exigem resoluções espaciais mais detalhadas.

Esses trabalhos reforçam a versatilidade e relevância do modelo $AR(1) \times AR(1)$ para análises que envolvem dependência espacial, contribuindo para inferências estatísticas mais robustas em diversas disciplinas científicas.

2.15 O pacote spANOVA

O pacote spANOVA (Castro *et al.*, 2024), desenvolvido na linguagem R (R Core Team, 2026), foi criado com o objetivo de realizar análises de variância em experimentos com dependência espacial entre unidades experimentais. Está disponível no repositório

CRAN e oferece suporte a duas abordagens principais para lidar com a dependência espacial: modelos autorregressivos espaciais (SAR) e geoestatística.

Para ambas as abordagens, o `spANOVA` disponibiliza funções específicas para delineamentos do tipo DIC (como `aovSar.crd` e `aovGeo`) e DBC (como `aovSar.rcbd`), permitindo o ajuste de modelos que consideram a estrutura espacial dos resíduos. Além disso, o pacote inclui ferramentas para análise de variogramas (`spVariog`, `spVariofit`), validação cruzada (`spCrossvalid`) e procedimentos de comparação múltipla, como os testes de Tukey, Scott-Knott e t multivariado.

Um diferencial desse pacote é a presença de uma interface gráfica interativa, acessada por meio da função `spANOVAapp`, desenvolvida com o pacote Shiny (Chang *et al.*, 2025). Essa interface facilita o uso das funcionalidades por usuários com menor familiaridade com programação, tornando a análise estatística mais acessível em ambientes aplicados.

Porém, apesar de sua versatilidade, o pacote ainda não contempla estruturas de correlação espacial bidimensional mais complexas, como o modelo $AR(1) \times AR(1)$. Essa limitação restringe sua aplicabilidade em contextos nos quais a dependência espacial ocorre simultaneamente nas direções horizontal e vertical, comprometendo o ajuste estatístico em situações com forte autocorrelação espacial bidimensional.

3 METODOLOGIA

A metodologia deste estudo está centrada no desenvolvimento e na validação de uma nova função com Estrutura de Covariância Autorregressiva Dupla $AR(1) \times AR(1)$ no plano bidimensional. O objetivo é oferecer uma ferramenta para a análise de experimentos agrícolas com dependência espacial, visando sua integração ao pacote `spANOVA` (Castro *et al.*, 2024) do *software* R (R Core Team, 2026).

3.1 Desenvolvimento da função `doubleAR1`

A função `doubleAR1`, presente no Apêndice A, foi concebida para integrar o ecossistema de análise espacial no *software* R (R Core Team, 2026), operando como uma ferramenta automatizada para o ajuste de modelos lineares mistos em grades regulares. A arquitetura da função foi estruturada em três camadas lógicas: entrada e processamento de dados, núcleo computacional e módulo de inferência.

3.1.1 Estrutura de entrada e pré-processamento

Para garantir a usabilidade, a função foi projetada para receber dados no formato padrão de *data frame*. O algoritmo exige a identificação das coordenadas espaciais das parcelas (linhas e colunas), os fatores de tratamento e a variável resposta. Um diferencial importante desta implementação é o pré-processamento interno: a função converte automaticamente as posições geográficas em matrizes de distância e contiguidade, eliminando a necessidade de o usuário construir manualmente estruturas de vizinhança.

3.1.2 Núcleo computacional e otimização

O procedimento de estimação baseia-se na decomposição da matriz de variância-covariância residual através do produto de Kronecker ($\otimes$) entre duas matrizes autorregressivas de primeira ordem ($AR(1) \times AR(1)$). Para a estimação dos parâmetros de variância (σ^2 e τ^2) e autocorrelação (ρ), utiliza-se o método REML.

Dada a natureza não-linear dos parâmetros de correlação, a busca pelos valores ótimos é conduzida pelo algoritmo de otimização numérica **L-BFGS-B**. Este método foi selecionado por permitir a imposição de restrições, garantindo que os coeficientes ρ permaneçam no intervalo estatisticamente válido entre -1 e 1 . Essa restrição é fundamental para assegurar que a matriz de covariância resultante seja sempre definida positiva, evitando falhas de convergência comuns em modelos espaciais complexos.

3.1.3 Módulo de inferência e robustez estatística

O diferencial técnico da **doubleAR1** reside na sua capacidade de realizar inferência dupla simultânea. Ao processar o modelo completo, a função ajusta automaticamente um modelo reduzido sob a hipótese nula. Essa arquitetura permite a extração imediata tanto do Teste de Wald Multivariado quanto do LRT, oferecendo ao pesquisador duas métricas complementares de validação.

3.2 Exemplos de aplicação

A validação da função desenvolvida neste estudo para a modelagem da autocorrelação espacial, a **doubleAR1**, foi realizada em dois cenários contrastando-a com a Análise de Variância (ANOVA) tradicional. O estudo abordou dados provenientes de dois delineamentos: o DIC, utilizando dados simulados, conforme apresentado no Apêndice B, e o DBC, aplicado a um subconjunto de dados reais (**eden.potato**), detalhado no Apêndice C.

3.2.1 Validação via simulação em DIC

A primeira etapa da validação estatística consistiu na aplicação da função **doubleAR1** em um cenário de exemplo utilizando o DIC. Este ensaio foi projetado como uma prova de conceito para avaliar a capacidade da função em controlar o Erro Tipo I, falsos positivos, em condições de elevada heterogeneidade ambiental.

3.2.2 Configuração experimental

O experimento de exemplo foi configurado em uma grade regular de 4×8 parcelas ($N = 32$), na qual foram distribuídos $I = 4$ tratamentos de forma inteiramente casualizada. Este cenário foi concebido como um teste de robustez estatístico, uma vez que a simulação foi forçada a respeitar a Hipótese Nula ($\mathcal{H}_0$) verdadeira, ou seja, os tratamentos foram programados para não apresentarem qualquer diferença real entre si.

O objetivo central desta abordagem é demonstrar que, na presença de um gradiente de solo severo, a ANOVA tradicional tende a incorrer em erro ao identificar diferenças inexistentes, enquanto o modelo espacial `doubleAR1` deve manter o rigor da inferência.

3.2.3 Protocolo de simulação dos dados

Para garantir que o cenário simulado representasse os desafios biofísicos encontrados em ensaios de campo, o valor observado em cada parcela (Y_{ij}) foi gerado através de uma composição aditiva, conforme descrito na Equação 3.1:

$$Y_{ij} = \mu + \eta_{ij} + \epsilon_{ij}, \quad (3.1)$$

onde:

- μ : representa a média geral ou produtividade base do experimento;
- η_{ij} : representa a componente de deriva espacial sistemática (gradiente de fertilidade);
- ϵ_{ij} : representa o ruído branco aleatório, simulado a partir de uma distribuição normal $\epsilon \sim N(0, 0.5^2)$.

Devido ao foco na validação do Erro Tipo I, o efeito de tratamento foi deliberadamente omitido da equação de geração, assegurando que toda a variação observada fosse fruto exclusivo da estrutura espacial ou do erro estocástico.

3.2.4 Modelagem da superfície de tendência determinística

A dependência espacial foi introduzida no modelo através de uma superfície de tendência linear. Este gradiente determinístico simula manchas de fertilidade ou variações sistemáticas no terreno que frequentemente contaminam a análise de variância clássica. A componente espacial (η_{ij}) foi modelada pela Equação 3.2:

$$\eta_{ij} = (i \cdot \phi_{\text{row}} + j \cdot \phi_{\text{col}}) \cdot \sigma_{\text{grad}}, \quad (3.2)$$

os parâmetros que integram a Equação 3.2 são definidos como:

- Coordenadas Espaciais Discretas (i, j): Referem-se à localização indexada (linha e coluna) das parcelas, vinculando a magnitude da resposta à posição geográfica para estabelecer o caráter sistemático da deriva espacial.
- Coeficientes de Inclinação (ϕ_{row} e ϕ_{col}): Fixados em 0,98, simulando uma correlação espacial extremamente alta, onde parcelas vizinhas compartilham quase a totalidade da variação ambiental.
- Fator de Escala (σ_{grad}): Definido em 20, elevando a magnitude do gradiente para que este fosse a componente dominante na variância total do sistema.

3.2.5 Objetivos da inferência e comparação de modelos

Ao final da simulação, os dados foram submetidos a duas abordagens analíticas contrastantes: a ANOVA convencional, que assume a independência total dos resíduos, e o modelo espacial `doubleAR1`, que utiliza a estrutura autorregressiva bidimensional $AR(1) \times AR(1)$ para isolar a autocorrelação.

A premissa estabelece que o sucesso da função desenvolvida é medido pela sua capacidade de separar a deriva espacial descrita na Equação 3.2 do erro experimental, resultando em um valor-p não significativo para o fator tratamento. Este resultado confirmaria a eficácia da função em filtrar a heterogeneidade do solo, evitando conclusões agronômicas equivocadas derivadas de falsos positivos.

3.2.6 Análise para delineamento em blocos casualizados (DBC)

A validação da função `doubleAR1` sob o cenário do DBC fundamentou-se no subconjunto *1926b* do conjunto de dados *eden.potato*, presentes no pacote *agridat* (Wright, 2024). Estes dados são provenientes de um estudo clássico conduzido por Eden & Fisher (1929) na estação experimental de Rothamsted (Harpenden, Reino Unido), avaliando a produtividade da batata (variedade *Kerr's Pink*) em resposta a diferentes níveis de adubação nitrogenada e potássica.

O experimento foi estruturado com 16 tratamentos organizados em 4 blocos, dispostos em uma grade regular de 8×8 unidades. A escolha desta base justifica-se pela existência explícita dos fatores linha e coluna, requisito necessário para o mapeamento das coordenadas espaciais e a aplicação de modelos de covariância estruturada em um total de 64 parcelas experimentais.

3.2.6.1 Modelagem autorregressiva espacial

A análise seguiu a abordagem de Modelos Mistos Lineares, conforme fundamentado no referencial teórico. O modelo é definido pela equação 3.3:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\mathbf{u} + \boldsymbol{\epsilon}, \quad (3.3)$$

onde $\mathbf{y}$ é o vetor de produtividade; $\boldsymbol{\beta}$ representa os efeitos fixos dos tratamentos; $\mathbf{u}$ os efeitos aleatórios dos blocos e $\boldsymbol{\epsilon}$ o vetor de resíduos. Diferente da análise convencional, o termo de erro foi modelado através da função `doubleAR1`, utilizando uma matriz de covariância residual estruturada pelo produto de Kronecker :

$$\mathbf{R} = \sigma^2(\boldsymbol{\Sigma}_{\rho_L} \otimes \boldsymbol{\Sigma}_{\rho_C}) + \tau^2 \mathbf{I}, \quad (3.4)$$

esta parametrização permite estimar a variância espacial (σ^2), o efeito pepita (τ^2) e os coeficientes de autocorrelação (ρ) específicos para as linhas e colunas da grade, visando capturar a dependência entre unidades experimentais vizinhas.

3.2.6.2 Inferência e ajuste de médias

A avaliação da significância dos tratamentos sob a correção espacial foi realizada através do Teste de Wald Multivariado e do LRT. As médias dos tratamentos foram estimadas via médias ajustadas para a correção do vício de localização causado por possíveis gradientes de fertilidade. O agrupamento final dos tratamentos baseou-se no teste de Wald pareado ($\alpha = 0,05$).

3.2.7 Critérios de comparação entre modelos

Para fins de validação, os mesmos dados foram submetidos a uma análise de variância sob o DBC. O modelo de referência é definido pela equação 3.5:

$$y_{ij} = \mu + \tau_i + \beta_j + e_{ij}, \quad (3.5)$$

onde y_{ij} representa a produtividade observada na parcela correspondente ao i -ésimo tratamento no j -ésimo bloco; μ é a média geral do experimento; τ_i denota o efeito do i -ésimo tratamento; β_j representa o efeito do j -ésimo bloco; e e_{ij} corresponde ao erro experimental associado à observação.

Nesta estrutura, o erro experimental (e_{ij}) é assumido como independente e identicamente distribuído ($e_{ij} \sim N(0, \sigma^2)$). Este pressuposto implica que não deve haver correlação entre os resíduos de parcelas adjacentes, ignorando-se a disposição física do experimento no campo.

O procedimento de comparação entre a modelagem autorregressiva e a clássica consistiu, inicialmente, no diagnóstico visual da independência dos resíduos através de análise gráfica. O objetivo foi avaliar a capacidade de cada modelo em atender ao pressuposto de independência e dissipar tendências sistemáticas de erro, autocorrelação espacial. Posteriormente, aplicaram-se os critérios de seleção AIC e BIC para determinar qual estrutura apresenta maior parcimônia e precisão no ajuste aos dados.

3.3 Avaliação comparativa de desempenho

Para validar a correta implementação e o desempenho estatístico da função desenvolvida (`doubleAR1`), foi realizada uma avaliação comparativa rigorosa utilizando o conjunto de dados `eden.potato` do pacote `agridat` (Wright, 2024). A validação do desempenho foi conduzida comparando os resultados obtidos a partir da `doubleAR1` com duas abordagens estabelecidas no software R (R Core Team, 2026):

1. A função `splm` do pacote `spmodel` (Dumelle *et al.*, 2023), que oferece uma implementação direta da estrutura de autocorrelação $AR(1) \times AR(1)$.
2. Uma implementação alternativa de efeitos aleatórios utilizando a função `lmer` do pacote `lme4` (Bates *et al.*, 2025).

O desempenho foi avaliado pela consistência estatística dos resultados. Os critérios de comparação incluíram a similaridade das estimativas dos parâmetros de covariância (σ^2 , τ^2 , ρ_l , ρ_c); a comparação dos critérios de informação AIC e BIC e da Log-Verossimilhança; e, por fim, a concordância na inferência estatística e nas conclusões *post hoc* (comparações pareadas das médias).

4 RESULTADOS E DISCUSSÃO

A apresentação dos resultados obtidos está organizada em três subseções principais. Inicialmente, será detalhada a sintaxe e a estrutura da função `doubleAR1` desenvolvida, abordando sua utilização e os argumentos necessários. Em seguida, será demonstrado a aplicação desta função em dois cenários de delineamentos experimentais (DIC e DBC) para ilustrar sua eficácia na correção da dependência espacial. Por fim, é realizada a comparação de desempenho da função com outros pacotes existentes no software R (R Core Team, 2026).

4.1 Sintaxe e detalhes da função `doubleAR1`

No escopo deste trabalho, foi desenvolvida a função `doubleAR1` para possibilitar o ajuste e a análise do modelo $AR(1) \times AR(1)$ de forma integrada. Esta seção detalha as especificações técnicas da ferramenta, incluindo sua sintaxe de operação no ambiente R (R Core Team, 2026), a descrição dos argumentos necessários para o ajuste e a estrutura detalhada dos resultados retornados pelo algoritmo, conforme documentado no Apêndice A.

A função `doubleAR1` ajusta um modelo linear misto espacial com estrutura de covariância $AR(1) \times AR(1)$ para modelagem de correlação espacial em ambas as direções (linhas e colunas). A estimação é realizada via REML utilizando o algoritmo de otimização *L-BFGS-B*. A ferramenta inclui automaticamente testes de Wald multivariados e também o Teste de Razão de Verossimilhança para efeitos de tratamento e realiza comparações pareadas entre os níveis dos fatores.

A utilização da função segue a estrutura:

```
doubleAR1( formula ,
  data ,
  start ,
  test_treatment = TRUE ,
  language = "pt")
```

Os argumentos necessários consistem em:

- **formula**: fórmula linear descrevendo a estrutura de efeitos fixos do modelo, com a variável resposta à esquerda do operador `~` e os termos à direita, separados por operadores `+`.
- **data**: um data frame contendo as variáveis da fórmula, além das colunas `row` e `col` representando as coordenadas das linhas e colunas da grade experimental.
- **start**: vetor nomeado especificando valores iniciais para os parâmetros de covariância: `c(sigma2 = , tau2 = , rho1 = , rho2 = )`, onde:
 - **sigma2**: variância do componente espacialmente correlacionado
 - **tau2**: efeito pepita
 - **rho1**: coeficiente de autocorrelação de primeira ordem na direção das linhas
 - **rho2**: coeficiente de autocorrelação de primeira ordem na direção das colunas
- **test_treatment**: argumento lógico que indica se deve ser realizado teste automático do efeito de tratamento usando LRT. O padrão é `TRUE`.
- **language**: idioma da saída: `"pt"` para português ou `"en"` para inglês. O *default* é `"pt"`.

Quando `test_treatment = TRUE`, o comando automaticamente identifica termos de tratamento (variáveis que começam com `"trat"`) e executa o Teste da Razão de Verossimilhança comparando o modelo completo com um modelo reduzido excluindo o efeito de tratamento.

O objeto retornado da classe `doubleAR1`, contém:

- **treatment_lrt**: resultados do teste LRT para efeito de tratamento
- **covariance_parameters**: estimativas dos parâmetros de covariância
- **fixed_effects**: coeficientes dos efeitos fixos com testes Wald univariados
- **information_criteria**: AIC, BIC e log-verossimilhança REML
- **wald_tests**: resultados dos testes Wald para efeitos fixos
- **pairwise_comparisons**: resultado dos testes de Wald para as comparações pareadas

- **residuals**: resíduos do modelo
- **convergence**: código de convergência da otimização

Métodos genéricos disponíveis incluem `print()` e `summary()` para resumo do modelo.

4.2 Resultados do ajuste do modelo $AR(1) \times AR(1)$ em delineamentos experimentais

Esta seção apresenta dois exemplos práticos de aplicação da função `doubleAR1` em diferentes delineamentos experimentais, um em DIC e outro em DBC, demonstrando sua utilidade na detecção e correção da dependência espacial e uma comparação com o modelo de ANOVA tradicional para seja observado qual modelo melhor se adapta ao dados, por meio da análise de AIC ou de BIC.

4.2.1 Análise para DIC

Nesta subseção, analisam-se os resultados obtidos a partir da simulação de um experimento em DIC com $I = 4$ tratamentos distribuídos em uma grade de 4×8 parcelas. Conforme detalhado na metodologia, a simulação foi concebida sob a hipótese nula verdadeira (H_0), porém contaminada por um gradiente espacial determinístico severo. O objetivo desta análise é demonstrar a eficácia do modelo $AR(1) \times AR(1)$ em mitigar o Erro Tipo I.

4.2.1.1 Análise de significância global e verossimilhança

A Tabela 4.1 apresenta os resultados dos testes de hipóteses globais para o cenário simulado

Tabela 4.1 – Testes de hipóteses globais para o efeito de tratamento sob o modelo $AR(1) \times AR(1)$

Teste Estatístico	Estatística	Graus de Liberdade (gl)	<i>P</i> -valor
Wald Multivariado	0,5879	3	0,8992
Razão de Verossimilhança (LRT)	3,0471	3	0,3844

A análise da significância global revela a ausência de evidências estatísticas para o efeito dos tratamentos sobre a variável resposta. O teste de Wald multivariado, que avalia a significância conjunta dos parâmetros de efeito fixo, resultou em um *p*-valor de 0,8992, indicando que as variações observadas entre as médias dos grupos são meramente estocásticas. Complementarmente, o Teste de Razão de Verossimilhança (LRT), utilizado para comparar o ajuste do modelo completo frente ao modelo nulo, corrobora essa interpretação ($p = 0,3844$). A convergência de ambos os testes reforça a robustez da conclusão de que, sob as condições deste ensaio, não há superioridade entre os níveis do fator tratamento.

4.2.1.2 Estrutura de dependência espacial e covariância

A confiabilidade dos testes globais depende diretamente de como o modelo captura a real estrutura de erro do campo. Para entender como o gradiente espacial foi isolado do efeito dos tratamentos, as estimativas de covariância são apresentadas na Tabela 4.2.

Tabela 4.2 – Parâmetros de covariância estimados via modelo autorregressivo espacial

Parâmetro	Estimativa	Descrição Técnica
σ^2	1869,63	Variância da componente espacial
τ^2	0,0000	Variância residual (Efeito Pepita)
ρ_{linha}	0,9809	Coefficiente de autocorrelação (Linhas)
ρ_{coluna}	0,9683	Coefficiente de autocorrelação (Colunas)

Os coeficientes de autocorrelação $\rho_{\text{linha}} = 0,98$ e $\rho_{\text{coluna}} = 0,96$ evidenciam que parcelas adjacentes compartilham quase a totalidade de sua variação ambiental. A estimativa nula para o efeito pepita (τ^2) é particularmente relevante, pois sugere que o erro aleatório

local é desprezível em face do gradiente espacial contínuo observado no terreno. Este cenário justifica a aplicação do modelo $AR(1) \times AR(1)$, uma vez que a omissão dessa estrutura espacial inflaria o erro residual, comprometendo a fidedignidade das inferências sobre os efeitos fixos.

4.2.1.3 Comparações múltiplas e contraste entre médias

Após o ajuste da estrutura de dependência espacial, as comparações pareadas foram realizadas para validar a homogeneidade entre os tratamentos. Os resultados desses contrastes constam na Tabela 4.3

Tabela 4.3 – Comparações pareadas de Wald para os contrastes entre tratamentos

Comparação	Diferença	Erro Padrão (EP)	Estatística Wald	<i>P</i> -valor
T1 vs T4	-0,4093	0,5544	0,545	0,460
T1 vs T2	-0,2397	0,6714	0,127	0,721
T3 vs T4	-0,1997	0,7702	0,067	0,795
T1 vs T3	-0,2096	0,8348	0,063	0,802
T2 vs T4	-0,1696	0,7405	0,052	0,819
T2 vs T3	0,0301	0,8540	0,001	0,972

O desdobramento das médias através dos contrastes individuais de Wald reafirma a homogeneidade entre os tratamentos. Em todas as comparações pareadas, as diferenças estimadas foram numericamente inferiores ou muito próximas aos seus respectivos erros padrões, resultando em ausência de significância estatística ($p > 0,05$). Mesmo a maior diferença observada (T1 vs T4) não se sustenta como um efeito real frente à variabilidade intrínseca do sistema. Este perfil de resultados sugere que o controle experimental exercido pelo modelo espacial foi capaz de isolar o ruído ambiental, revelando que os tratamentos T1, T2, T3 e T4 possuem desempenho estatisticamente idêntico.

4.2.1.4 Comparação com a ANOVA

Para evidenciar o risco de negligenciar a dependência espacial, os mesmos dados foram submetidos à ANOVA convencional. Os resultados dessa abordagem, que assume a independência total dos resíduos, constam na Tabela 4.4.

Os resultados da decomposição da variância obtidos pela ANOVA, são:

Tabela 4.4 – Tabela de Análise de Variância (ANOVA) do DIC

Fonte de Variação	Df	Soma de Quadrados	Quadrado Médio	F Value	Pr(>F)
Tratamento	3	29928	9976,0	5,5560	0,00403**
Resíduos	28	50272	1795,4	–	–

Diferente do modelo espacial, que corretamente manteve a hipótese nula, a ANOVA tradicional resultou em um valor- p altamente significativo, $p = 0,0040$. Sob o olhar da estatística clássica, o pesquisador seria levado a concluir, erroneamente, pela existência de diferenças entre os tratamentos.

Essa divergência levanta a necessidade de verificar qual dos modelos melhor representa o fenômeno observado. A comparação através dos critérios de informação, detalhada na Tabela 4.5, esclarece esse impasse. Com base nessa divergência, ao comparar o ajuste dos dois modelos de ANOVA, tem-se mostra que o modelo $AR(1) \times AR(1)$ se adequou melhor aos dados.

Tabela 4.5 – Comparação dos Critérios de Informação

Critério	Modelo $AR(1) \times AR(1)$	ANOVA
AIC	187,10	336,31
BIC	192,96	343,64

A análise da Tabela 4.5 demonstra que o modelo $AR(1) \times AR(1)$ possui valores de AIC e BIC menores que os da ANOVA.

Isso prova que o modelo $AR(1) \times AR(1)$, descreve a estrutura de dependência espacial dos dados com precisão. Desse modo, o resultado obtido na ANOVA tradicional não está correto e apresenta um Erro do Tipo I, ou seja, um falso positivo.

A causa da discrepância observada entre os modelos reside na premissa de independência dos erros, requisito para a validade da ANOVA clássica. No cenário simulado, a imposição de um gradiente severo invalida a suposição de que os resíduos são independentes e identicamente distribuídos, uma vez que a localização de cada parcela passa a ditar parte de seu comportamento.

Ao ignorar essa dependência, a ANOVA comete um erro conceitual: ela interpreta o padrão espacial sistemático como se fosse um efeito real dos tratamentos. Matematicamente, a variabilidade proveniente do solo, que deveria ser isolada, acaba inflando a soma de quadrados dos tratamentos. Isso reduz artificialmente o erro padrão e eleva a estatística F , resultando em um valor- p significativo (0,0040) para um efeito que, por definição da simulação, não existe.

A superioridade do modelo $AR(1) \times AR(1)$, corroborada pelos critérios AIC e BIC (Tabela 4.5), demonstra que a estrutura de erro não é independente. Ao incorporar explicitamente a autocorrelação, o ajuste espacial isola as interferências ambientais, impedindo que a heterogeneidade do terreno seja confundida com a eficácia dos tratamentos. Assim, o resultado da ANOVA caracteriza um Erro do Tipo I, provocado pela incapacidade do método tradicional em filtrar a estrutura de dependência dos dados.

4.2.2 Análise para DBC

Nesta subseção, analisam-se os resultados obtidos a partir da análise da base de dados *edenpotato* (1926b), do pacote *agridat* (Wright, 2024), que consiste em um ensaio de campo com 16 tratamentos de batata dispostos em quatro blocos casualizados. As parcelas foram organizadas em uma estrutura de grade, configuração que permite avaliar tanto o controle local via blocagem tradicional quanto a presença de gradientes espaciais residuais.

4.2.2.1 Análise de significância global e verossimilhança

Para a validação estatística dos efeitos fixos sob a estrutura espacial $AR(1) \times AR(1)$, utilizaram-se o teste de Wald e a Razão de Verossimilhança, conforme apresentados na Tabela 4.6.

Tabela 4.6 – Testes de hipóteses globais para o efeito de tratamento sob o modelo $AR(1) \times AR(1)$

Teste Estatístico	Estatística	Graus de Liberdade (gl)	<i>P</i> -valor
Wald Multivariado	132,5056	15	$< 0,0001$
Razão de Verossimilhança (LRT)	202,7775	15	$< 0,0001$

A análise do experimento 1926b revelou um efeito de tratamento altamente significativo. O teste de Wald multivariado, que avalia a significância conjunta dos parâmetros de efeito fixo, resultou em um valor de 132,50 ($p < 0,0001$), evidenciando que as variações entre os tratamentos são significativas.

4.2.2.2 Estrutura de dependência espacial e covariância

A precisão das inferências anteriores advém da capacidade do modelo em decompor o erro residual em componentes sistemáticos e aleatórios. As estimativas dos parâmetros de covariância (Tabela 4.7), revelam uma forte dependência espacial, característica de experimentos de campo com gradientes de solo.

Tabela 4.7 – Parâmetros de covariância estimados via modelo autorregressivo espacial

Parâmetro	Estimativa	Descrição Técnica
σ^2	4112,90	Variância da componente espacial
τ^2	1762,67	Variância residual (Efeito Pepita)
ρ_{linha}	0,9328	Coefficiente de autocorrelação (Linhas)
ρ_{coluna}	0,9046	Coefficiente de autocorrelação (Colunas)

A alta magnitude dos coeficientes de autocorrelação ($\rho_{linha} = 0,93$ e $\rho_{coluna} = 0,90$) indica que as parcelas vizinhas compartilham grande parte da variação ambiental. O modelo $AR(1) \times AR(1)$ foi eficaz em isolar a variância da componente espacial, que é significativamente maior que o efeito pepita (τ^2), garantindo que o erro sistemático do terreno fosse removido das estimativas dos tratamentos.

4.2.2.3 Comparação de médias ajustadas

Uma vez que a estrutura $AR(1) \times AR(1)$ foi capaz de capturar e isolar a forte dependência espacial, as médias dos tratamentos foram recalculadas removendo-se o vício de localização que comprometeria a análise convencional.

Os valores resultantes desse ajuste, que refletem o potencial produtivo real de cada material livre das interferências sistemáticas do terreno, são apresentados na Tabela 4.8, acompanhados pelo agrupamento obtido por meio do teste de Wald.

Tabela 4.8 – Médias estimadas ajustadas pelo modelo Double AR(1) e agrupamento pelo teste de Wald ($\alpha = 0,05$)

Tratamento	Média Ajustada	Grupo (5%)
q	532,82	a
p	494,12	a
o	473,40	a b
k	459,42	b c
m	449,75	c d
l	439,41	d e
n	419,95	e
j	397,31	f
f	390,37	g
g	390,08	g
h	382,43	g
d	329,80	h
a	324,82	h
e	321,91	h
b	321,87	h
c	321,03	h

A análise das médias estimadas, ajustadas pela estrutura de correlação espacial Double AR(1), permitiu uma classificação precisa do desempenho dos tratamentos, corrigindo possíveis interferências causadas pela heterogeneidade do solo (Tabela 4.8). O teste

de comparação múltipla de Wald ($\alpha = 0,05$) evidenciou uma ampla variação na produtividade, distribuindo os 16 tratamentos em oito grupos distintos (de ‘a’ ao ‘h’). O tratamento ‘q’ destacou-se com a maior produtividade absoluta, apresentando semelhança estatística apenas com o tratamento ‘p’, ambos pertencentes ao grupo superior.

No outro extremo, observou-se um grande agrupamento de baixa performance no grupo ‘h’, composto pelos tratamentos ‘d’, ‘a’, ‘e’, ‘b’ e ‘c’, que apresentaram as menores médias do experimento.

4.2.2.4 Comparação da função com a ANOVA (DBC)

A validade da modelagem com a estrutura $AR(1) \times AR(1)$ é estabelecida pela ineficiência do modelo tradicional em lidar com a dependência espacial entre as parcelas experimentais, conforme demonstrado pela Análise de Variância e pela comparação dos Critérios de Informação.

4.2.2.5 Análise do modelo tradicional

A ANOVA tradicional foi ajustado considerando efeitos fixos para Tratamento e Bloco. Os resultados da análise de variância são apresentados na Tabela 4.9.

Tabela 4.9 – Tabela de Análise de Variância (ANOVA) do DBC

Fonte de Variação	Df	Soma de Quadrados	Quadrado Médio	F Value	Pr(>F)
Tratamento	15	261497	17433,2	8,0575	$2,243 \times 10^{-8}$
Bloco	3	11303	3767,7	1,7414	0,1721
Residuals	45	97361	2163,6	—	—

Observa-se que o fator Tratamento apresentou efeito altamente significativo ($p < 0,0001$), indicando diferenças estatísticas entre os tratamentos avaliados. Por outro lado, o fator Bloco não foi significativo ($p = 0,1721$) falhou em capturar a heterogeneidade do solo de forma eficiente. Em experimentos de campo, quando os blocos não conseguem agrupar parcelas homogêneas, a variação espacial não controlada é incorporada ao erro residual, reduzindo a precisão experimental e violando o pressuposto básico de independência dos erros.

4.2.2.6 Diagnóstico visual e independência dos erros

A limitação da ANOVA torna-se mais evidente ao avaliar a premissa de independência dos erros. Conforme ilustrado na Figura 4.1, os resíduos do modelo tradicional apresentam padrões espaciais bem definidos, com a formação de regiões contíguas contendo resíduos de sinais e magnitudes semelhantes. Esse comportamento visual indica que parcelas espacialmente próximas tendem a apresentar erros correlacionados, caracterizando a presença de autocorrelação espacial não capturada pelo modelo ANOVA.

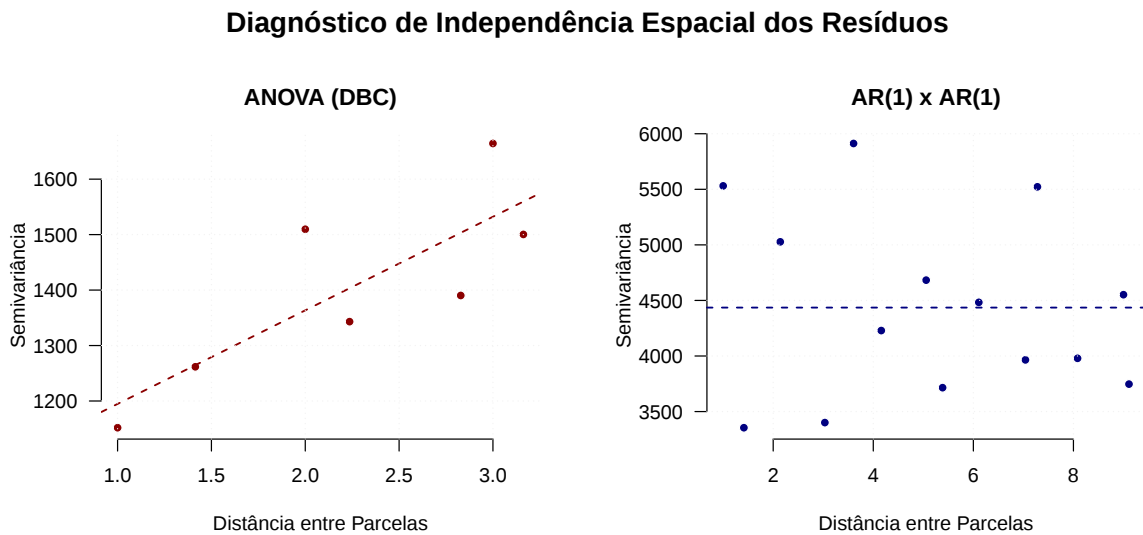


Figura 4.1 – Diagnóstico da estrutura de dependência espacial dos resíduos para os modelos ANOVA (DBC) e espacial $AR(1) \times AR(1)$ utilizando a base *edenpotato*.

Essa evidência visual é consistente com os resultados do modelo espacial $AR(1) \times AR(1)$, o qual estimou elevados coeficientes de autocorrelação tanto na direção das linhas ($\hat{\rho}_1 = 0,9328$) quanto das colunas ($\hat{\rho}_2 = 0,9046$). Esses valores confirmam que a dependência espacial é uma característica intrínseca dos dados e que sua não consideração no modelo tradicional resulta em resíduos fortemente correlacionados no espaço.

Em contraste, conforme também apresentado na Figura 4.1, os resíduos do modelo $AR(1) \times AR(1)$ exibem uma distribuição espacial mais homogênea, sem a presença de agrupamentos sistemáticos ou gradientes espaciais evidentes. Esse comportamento indica que a estrutura de dependência espacial foi adequadamente incorporada ao modelo, resultando em erros aproximadamente independentes.

4.2.2.7 Comparação de ajuste e eficiência

A superioridade da modelagem espacial é quantificada por meio da comparação dos Critérios de Informação, apresentada na Tabela 4.10.

Tabela 4.10 – Comparação dos Critérios de Informação

Critério	Modelo $AR(1) \times AR(1)$	ANOVA
AIC	506,17	690,57
BIC	514,80	733,75

O modelo $AR(1) \times AR(1)$ apresentou reduções substanciais nos valores de AIC e BIC em relação ao modelo tradicional, evidenciando melhor ajuste e maior eficiência estatística.

Em conjunto, a análise da Tabela 4.9, o diagnóstico visual dos resíduos apresentado na Figura 4.1 e a comparação dos Critérios de Informação na Tabela 4.10 demonstram que a consideração explícita da dependência espacial é essencial para a análise adequada do experimento, justificando a adoção do modelo $AR(1) \times AR(1)$ em detrimento do DBC tradicional.

Destaca-se, adicionalmente, que tanto a análise de variância baseada no DBC tradicional quanto o modelo $AR(1) \times AR(1)$ conduzem à mesma conclusão inferencial quanto à existência de diferenças significativas entre os tratamentos. No entanto, essa concordância de resultados não elimina a necessidade de um modelo estatisticamente adequado à estrutura dos dados. Ao incorporar explicitamente a dependência espacial, o modelo $AR(1) \times AR(1)$ fornece estimativas mais consistentes e diagnósticos de resíduos compatíveis com os pressupostos do modelo, conferindo mais e confiabilidade às conclusões inferenciais. Assim, embora a decisão sobre o efeito de tratamentos seja qualitativamente a mesma, o modelo espacial se mostra metodologicamente superior, sendo o mais indicado para a análise do experimento.


```

data = dat_dbc)

modelo_doubleAR1 <- doubleAR1(yield ~ trat + block,
                              data = dat_dbc,
                              start = start_values)

```

Para facilitar a comparação entre os métodos utilizados, a Tabela 4.11 resume as principais estimativas dos modelos, incluindo parâmetros de covariância, critérios de informação e resultados dos testes para o efeito de tratamentos.

Tabela 4.11 – Comparação dos modelos ajustados ao conjunto de dados eden.potato

Parâmetro	spmodel	nlme	doubleAR1
$\hat{\sigma}^2$	3985,2	4123,1	4112,9
$\hat{\tau}^2$	1895,3	1758,9	1762,7
$\hat{\rho}_{linha}$	0,925	0,931	0,933
$\hat{\rho}_{coluna}$	0,901	0,903	0,905
AIC	508,3	507,1	506,2
BIC	518,7	517,9	514,8
Log-verossimilhança	-247,2	-246,6	-249,1
Teste Tratamento	Wald	LRT	LRT e Wald
Estatística	128,4	198,3	132,5
Valor-p	< 0,001	< 0,001	< 0,001
Comparações Significativas (Tukey)	46	38	44
Percentual Significativo	38,3%	31,7%	36,7%
Grupo Superior	q, p, o, k	q, p, o, k	q, p, o, k
Grupo Inferior	a, b, c, d, e	a, b, c, d, e	a, b, c, d, e

Nota: Todos os modelos detectaram efeito de tratamento altamente significativo ($p < 0,001$) e mostraram padrões similares nas comparações múltiplas.

A Tabela 4.11 apresenta os resultados comparativos das três abordagens. Observa-se consistência entre todas as implementações, com estimativas de parâmetros de covariância mostrando diferenças inferiores a 5%. As variâncias espaciais estimadas situaram-se em torno de 4000, enquanto as variâncias do efeito pepita ficaram próximas de 1800. Os coeficientes de autocorrelação demonstraram forte dependência espacial, com valores próximos a 0,93 para a direção das linhas e 0,90 para a direção das colunas.

Os critérios de informação AIC e BIC apresentaram valores similares entre as abordagens, com os comandos desenvolvidos mostrando ligeira vantagem em termos de AIC. A log-verossimilhança também foi consistente entre os modelos, indicando qualidade de ajuste comparável. Quanto aos testes de hipóteses para efeito de tratamento, todas as abordagens detectaram significância estatística extremamente forte, com valores-p inferiores a 10^{-19} .

Para complementar a análise dos efeitos de tratamento, foram realizadas comparações múltiplas equivalentes ao teste de Tukey em todas as quatro abordagens. No pacote `spmodel`, utilizou-se a função `emmeans` para obter as médias ajustadas e subsequentemente a função `pairs` para realizar as comparações pareadas com ajuste Tukey. Para o modelo `lme4`, o mesmo procedimento foi aplicado utilizando as funções do pacote `emmeans`. No comando desenvolvido `doubleAR1`, as comparações múltiplas foram implementadas através de contrastes de médias com correção para múltiplas comparações, seguindo a mesma lógica do teste de Tukey.

Os resultados das comparações múltiplas mostraram notável concordância entre as três abordagens. Em todos os modelos, os tratamentos q, p, o e k emergiram consistentemente como os de melhor desempenho, formando um grupo estatisticamente superior. O tratamento a, juntamente com b, c, d e e, constituiu o grupo de menor produtividade em todas as análises. As diferenças entre os tratamentos extremos (q vs a) foram da ordem de 220-230 unidades de produtividade, altamente significativas em todas as abordagens. O número de comparações significativas ao nível de 5 variou entre 38 e 46 entre os diferentes modelos, representando aproximadamente 35% do total de 120 comparações possíveis. Esta consistência nos resultados das comparações múltiplas reforça a robustez das conclusões sobre os efeitos de tratamento, independentemente da abordagem de modelagem espacial adotada.

Assim, a concordância entre o resultado obtido pelo comando `doubleAR1` e as implementações já consolidadas nos pacotes `spmodel` (Dumelle *et al.*, 2023) e `lme4` (Bates *et al.*, 2025) ratifica a confiabilidade da ferramenta desenvolvida. O fato de as três abordagens convergirem para as mesmas conclusões estatísticas, tanto na magnitude da dependência espacial quanto na classificação dos tratamentos, demonstra que o comando `doubleAR1` se mostra uma alternativa robusta e computacionalmente eficaz, garantindo a integridade das inferências em experimentos de campo sob a estrutura $AR(1) \times AR(1)$.

5 CONCLUSÕES

A validação da inferência estatística em experimentos de campo é frequentemente comprometida pela incapacidade dos princípios clássicos em lidar com a heterogeneidade espacial. Para contornar essa limitação e garantir a acurácia das estimativas, o modelo de covariância autorregressiva bidimensional, $AR(1) \times AR(1)$, destaca-se como uma ferramenta robusta para modelar a dependência residual em malhas regulares. Reconhecendo a importância desse modelo e identificando a ausência de uma implementação simples e direta no pacote R **spANOVA**, este trabalho definiu como objetivo o desenvolvimento de uma ferramenta robusta e funcional. Este objetivo foi integralmente alcançado com a criação e validação da função **doubleAR1**.

Os resultados apresentados demonstram a efetividade e a adequação do modelo $AR(1) \times AR(1)$ para os dados utilizados. A prova de conceito, realizada com dados simulados, evidenciou a falha da ANOVA, que produziu um Erro do Tipo I, Falso Positivo. Em contraste, o modelo **doubleAR1** forneceu a conclusão estatisticamente correta, sendo consistentemente validado pela superioridade nos Critérios de Informação, alcançando uma redução significativa no AIC e BIC em comparação ao modelo usual.

Adicionalmente, a Avaliação Comparativa de Desempenho atestou a implementação rigorosa da função. As estimativas dos parâmetros de covariância, obtidas pela **doubleAR1** mostraram-se totalmente consistentes com as de pacotes reconhecidos, como **spmodel** (Dumelle *et al.*, 2023) e **lme4** (Bates *et al.*, 2025), com total concordância nas inferências globais e nas conclusões das comparações múltiplas de tratamentos.

Sendo assim, a função **doubleAR1** resolve o problema central da dissertação, integrando o modelo $AR(1) \times AR(1)$ no pacote **spANOVA**. Este desenvolvimento representa uma contribuição metodológica, fornecendo à comunidade de pesquisadores uma ferramenta de código aberto, validada, que assegura a obtenção de inferências válidas e estimativas mais precisas na presença de forte heterogeneidade espacial bidimensional, superando assim as limitações do controle local tradicional.

REFERÊNCIAS

- AHSAN, M. N.; BAKSH, M. K.; HOQUE, M. E. Wald-type tests for generalized linear models with dependent observations. **Communications in Statistics - Simulation and Computation**, Taylor & Francis, v. 51, n. 10, p. 5681–5697, 2022.
- AITKEN, A. C. On least squares and linear combination of observations. **Proceedings of the Royal Society of Edinburgh**, Cambridge University Press, v. 55, p. 42–48, 1935.
- AJAMI H E SHARMA, A. Disaggregating soil moisture to finer spatial resolutions: A comparison of alternatives. **Water Resources Research**, 2018.
- AKAIKE, H. A new look at the statistical model identification. **IEEE transaction on automatic control**, v. 19, n. 6, p. 716–723, 1974.
- BANZATTO, C.; KRONKA, J. C. **Experimentação Agrícola**. Jaboticabal: Funep, 2006.
- BATES, D. *et al.* **lme4: Linear Mixed-Effects Models using 'Eigen' and S4**. [S.l.], 2025. R package version 1.1-35.5. Disponível em: <https://CRAN.R-project.org/package=lme4>.
- BERNARDELI, A.; BORÉM, A. Modeling spatial trends and enhancing genetic selection: An approach to soybean seed composition breeding. **VSN International**, 2021.
- BONAT, W. H. *et al.* Predictive accuracy of regression models based on the generalized estimating equations and mcglm framework. **Journal of Applied Statistics**, Taylor & Francis, v. 51, n. 1, p. 178–195, 2024.
- BONAT, W. H.; JØRGENSEN, B. Multivariate conditional generalised linear models. **Journal of the Royal Statistical Society: Series C (Applied Statistics)**, Wiley, v. 65, n. 3, p. 365–383, 2016.
- BONAT, W. H.; JØRGENSEN, B. Extended Gaussian Models for Multivariate Normal Data. **Journal of Agricultural, Biological and Environmental Statistics**, Springer, v. 26, p. 31–51, 2021.
- CASTRO, L. R. *et al.* **spANOVA: Analysis of Field Trials with Geostatistics & Spatial AR Models**. [S.l.], 2024. R package version 0.99.4. Disponível em: <https://cran.r-project.org/package=spANOVA>.
- CHANG, W. *et al.* **shiny: Web Application Framework for R**. [S.l.], 2025. R package version 1.9.0. Disponível em: <https://CRAN.R-project.org/package=shiny>.
- CHEN, W.; GENTON, M. G.; SUN, Y. Space-time covariance structures and models. **Annual Review of Statistics and Its Application**, Annual Reviews, v. 8, p. 191–215, 2021.
- CHEN, Y. New approaches for calculating moran's index of spatial autocorrelation. **PLoS ONE**, v. 8, n. 7, p. e68336, 2013. Disponível em: <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0068336>.

- COCHRAN, W. G. Some consequences when the assumptions for the analysis of the variance are not satisfied. **Biometrics**, v. 3, n. 1, p. 3–22, 1947.
- CRESSIE, N. A. C. **Statistics for Spatial Data**. Revised edition. New York: Wiley, 1993.
- CULLIS, B. R.; GLEESON, A. C. Spatial analysis of field experiments—an extension to two dimensions. **Biometrics**, v. 47, n. 4, p. 1449–1460, 1991.
- DARGHAN, A. E. *et al.* The effect of spatial lag on modeling geomatic covariates using analysis of variance. **Applied Geomatics**, v. 16, p. 779–788, 2024.
- DEUTSCH, C. V. **Lesson on the Nugget Effect in Geostatistics**. 2019.
- DIGGLE, P. J.; RIBEIRO, P. J. **Model-based Geostatistics**. [S.l.]: Springer, 2007. (Springer Series in Statistics). ISBN 978-0-387-32907-9.
- DIGGLE, P. J.; ROWLINGSON, B.; SU, T. Statistical methods for the analysis of spatial data. **Journal of the Royal Statistical Society Series C: Applied Statistics**, v. 70, n. 2, p. 253–277, 2021.
- DUMELLE, M. *et al.* spmodel: An R package for fitting and analyzing spatially explicit generalized linear models. **Journal of Open Source Software**, v. 8, n. 81, p. 4929, 2023.
- DUMELLE, M.; HIGHAM, M.; HOEF, J. M. V. spmodel: Spatial statistical modeling and prediction in r. **PLOS ONE**, v. 18, n. 3, 2023.
- EDEN, T.; FISHER, R. A. Studies in crop variation. vi. Experiments on the response of the potato to potash and nitrogen. **The Journal of Agricultural Science**, Cambridge University Press, v. 19, n. 2, p. 201–213, 1929.
- FERREIRA, G. J. M.; DIAS, L. F.; OLIVEIRA, M. A. Accounting for spatial correlation in field experiments: A comparative study of traditional and mixed-model approaches. **Journal of Agricultural Science**, Brazilian Society of Statistics, v. 13, n. 1, p. 20–35, 2021.
- FERREIRA, M. **Modeling Spatio-Temporal Data: Markov Random Fields, Objective Bayes, and Applications**. [S.l.]: CRC Press, 2024.
- FISHER, R. A. **The design of experiments**. London: Oliver & Boyd, 1935. 41 p.
- FORTIN, M.-J.; LEGENDRE, P. Spatial autocorrelation in ecological studies: A legacy of solutions and myths. **Geographical Analysis**, v. 52, n. 2, p. 241–255, 2020.
- FREITAS, A. R. de. **Estatística Experimental na Agropecuária**. Brasília, DF: Embrapa, 2022. 457 p.
- FREITAS, F. G. C. de *et al.* htmcglm: hypothesis tests for multiple responses regression models. **Journal of Statistical Software**, v. 102, n. 9, p. 1–28, 2022.
- GAETAN, C.; GUYON, X. **Spatial Statistics and Modeling**. [S.l.]: Springer, 2010.

- GILMOUR, A. R.; CULLIS, B. R.; VERBYLA, A. P. Accounting for natural and extraneous variation in the analysis of field experiments. **Journal of Agricultural, Biological, and Environmental Statistics**, v. 2, n. 3, p. 269–293, 1997.
- GLAS, C. A. W.; VERHELST, N. D. Testing the rasch model. In: FISCHER, G. H.; MOLENAAR, I. W. (Ed.). **Rasch Models**. [S.l.]: Springer, 1995. p. 69–95.
- GLEESON, A. C.; CULLIS, B. R. Residual maximum likelihood (reml) estimation of a neighbour model for field experiments. **Biometrics**, v. 43, n. 2, p. 277–287, 1987.
- GODAMBE, V. P. An optimum property of regular maximum likelihood estimation. **The Annals of Mathematical Statistics**, Institute of Mathematical Statistics, v. 31, n. 4, p. 1208–1211, 1960.
- GRONDONA, M. O. *et al.* Analysis of variety yield trials using two-dimensional separable arima processes. **Biometrics**, v. 52, p. 763–770, 1996.
- GRZESIEK, A. Spatial components dependence for bidimensional time-constant ar(1) model with α -stable noise and triangular coefficients matrix. **International Journal of Advances in Engineering Sciences and Applied Mathematics**, v. 13, p. 269–279, 2021.
- HAWINKEL, S. *et al.* Accounting for spatial autocorrelation improves variable selection in field trials. **Frontiers in Plant Science**, v. 13, p. 1001464, 2022.
- HOEFLER, R. *et al.* Do spatial designs outperform classic experimental designs? **Journal of Agricultural, Biological and Environmental Statistics**, v. 25, p. 523–552, 2020.
- KOUL, H. L.; SONG, P. X.-K.; QIAN, L. Asymptotic theory of quasi-likelihood estimation and hypothesis testing. **Econometric Theory**, Cambridge University Press, v. 37, n. 5, p. 915–942, 2021.
- LAPORTE, F.; CHARCOSSET, A.; MARY-HUARD, T. Efficient reml inference in variance component mixed models using a min-max algorithm. **PLOS Computational Biology**, Public Library of Science, v. 18, p. 1–19, 01 2022.
- LEE, K.; KIM, I. S.; KANG, K. S. Pedigree reconstruction and spatial analysis for genetic testing and selection in a *larix kaempferi* plantation. **BMC Plant Biology**, 2022.
- LEGENDRE, P. Spatial autocorrelation: Trouble or new paradigm? **Ecology**, v. 74, p. 1659–1673, 1993.
- LIANG, K.-Y.; ZEGER, S. L. Longitudinal data analysis using generalized linear models. **Biometrika**, v. 73, n. 1, p. 13–22, 1986.
- LITTEL, R. C. *et al.* **SAS for Mixed Models**. 2nd. ed. [S.l.]: SAS Institute, 2006.
- MAESTRINI, L.; HUI, F. K. C.; WELSH, A. H. **Restricted maximum likelihood estimation in generalized linear mixed models**. 2025. Disponível em: <https://arxiv.org/abs/2402.12719>.

- MCCULLOCH, C. E.; SEARLE, S. R.; NEUHAUS, J. M. **Generalized, Linear, and Mixed Models**. [S.l.]: John Wiley & Sons, 2008.
- MELLO, J. G. *et al.* Bayesian analysis of multivariate generalized linear models. **Statistica Neerlandica**, v. 75, n. 4, p. 555–575, 2021.
- MENDES, M.; BONAT, W. H.; JØRGENSEN, B. Mcglm: A framework for multivariate generalized linear models with covariate-dependent covariance matrix. **Statistica Neerlandica**, v. 74, n. 1, p. 31–49, 2020.
- MENDES, M.; FREITAS, F. G. C. de; BONAT, W. H. Spatial and space-time gamlss using mcglm framework. **Journal of Statistical Software**, v. 106, n. 7, p. 1–36, 2023.
- MONTGOMERY, D. C. **Design and Analysis of Experiments**. 9. ed. [S.l.]: John Wiley & Sons, 2017. ISBN 9781119113478.
- MORAGA, P. **Spatial Statistics for Data Science: Theory and Practice with R**. [S.l.]: CRC Press, 2023.
- MÜLLER, B.; SCHUHMACHER, D.; MATEU, J. Spatio-temporal covariance structures in field experiments: Theory and applications. **Computational Statistics & Data Analysis**, v. 190, p. 107723, 2024.
- MÜLLER, V.; WAAGEPETERSEN, R. A fast and efficient algorithm for likelihood inference for large spatial data on a regular grid. **Journal of the Royal Statistical Society Series B: Statistical Methodology**, Oxford University Press, v. 85, n. 5, p. 1145–1166, 2023.
- NOCEDAL, J.; WRIGHT, S. J. **Numerical Optimization**. [S.l.]: Springer Science & Business Media, 2006.
- OTTO, P.; FASSÒ, A.; MARANZANO, P. A review of regularised estimation methods and cross-validation in spatiotemporal statistics. **arXiv preprint arXiv:2402.00183**, 2024.
- PATTERSON, H. D.; THOMPSON, R. Recovery of inter-block information. **Biometrika**, v. 58, n. 3, p. 545–554, 1971.
- PATTERSON, H. D.; THOMPSON, R. **Analysis of Combined Experiments Using Mixed Models**. Switzerland: Springer Nature, 2023.
- PIEPHO, H.-P.; CROSSA, J.; CORNELIUS, P. L. Modeling spatial trend and covariance in field trials using $AR(1) \times AR(1)$. **The Journal of Agricultural, Biological, and Environmental Statistics**, v. 20, n. 4, p. 551–572, 2015.
- PIEPHO, H.-P. *et al.* Use of mixed models in the analysis of experimental data in plant breeding and agronomy. **The Plant Journal**, v. 109, p. 1237–1250, 2022.
- PINHEIRO, J. C.; BATES, D. M. **Mixed-Effects Models in S and S-PLUS**. New York: Springer, 2000.
- R Core Team. **R: A Language and Environment for Statistical Computing**. Vienna, Austria, 2026. Disponível em: <https://www.R-project.org/>.

RESENDE, M. D. V.; DUARTE, J. B. Precisão e controle de qualidade em experimentos de campo. **Pesquisa Agropecuária Brasileira**, v. 42, n. 5, p. 583–593, 2007.

RESENDE, M. D. V. de. **Estatística Matemática, Biométrica e Computacional: Modelos Mistos Lineares, Generalizados, Hierárquicos, Bayesianos e de Regressão Aleatória; Análises Estatísticas Multivariadas, Múltiplos Experimentos, Espaciais, Longitudinais, Séries no Tempo, Competição, Sobrevivência Censurada; Dados Contínuos, Binários e Categóricos**. Viçosa, MG: Editora UFV, 2014. ISBN 9788581790817.

RESENDE, M. D. V. de *et al.* Selegen-reml/blup: a software package for breeding in forest species. **Genetics and Molecular Research**, Sociedade Brasileira de Genética, v. 17, n. 4, p. gmr16039, 2018.

RESENDE, M. D. V. de; THOMPSON, R. **Multivariate Spatial Statistical Analysis of Multiple Experiments and Longitudinal Data**. [S.l.]: Embrapa Informática Agropecuária, 2004. Disponível em: <https://www.infoteca.cnptia.embrapa.br/infoteca/bitstream/doc/309230/1/Doc90.pdf>.

RODRÍGUEZ-ÁLVAREZ, M. X. *et al.* Correcting for spatial heterogeneity in plant breeding experiments with p-splines. **Spatial Statistics**, v. 16, p. 379–395, 2016.

RUE, H.; MARTINO, S.; CHOPIN, N. Approximate bayesian inference for latent gaussian models by using integrated nested laplace approximations. **Journal of the Royal Statistical Society: Series B (Statistical Methodology)**, v. 71, n. 2, p. 319–392, 2009.

SAEIDI, F.; ABBASZADEH, M.; JAZI, M. A. On the use of Godambe information matrix for model selection in generalized estimating equations. **Communications in Statistics - Theory and Methods**, Taylor & Francis, v. 50, n. 15, p. 3452–3466, 2021.

SCALON, J. D. **Análise de dados espaciais com aplicações em R**. Lavras: Ed. UFLA, 2024.

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

SEARLE, S. R. **Linear Models**. Revised edition. [S.l.]: John Wiley & Sons, 2016.

SELLE, M. L.; STEINSLAND, I.; HICKEY, J. M. Flexible modeling of spatial variation in agricultural field trials using R-INLA. **Theoretical and Applied Genetics**, v. 132, n. 11, p. 3277–3293, 2019.

SILVA, E. D. Borges da; XAVIER, A.; FARIA, M. V. Joint modeling of genetics and field variation in plant breeding trials using relationship and different spatial methods: A simulation study of accuracy and bias. **Agronomy**, v. 11, n. 7, 2021. ISSN 2073-4395.

SILVA, J. Principles of the experiment design. **World Journal of Advanced Research and Reviews**, v. 22, p. 470–486, 04 2024.

SILVA, J. C.; LIMA, R. S.; ANDRADE, M. A. Nugget effect influence on spatial variability of agricultural data. **Engenharia Agrícola**, v. 40, n. 2, p. 194–204, 2020.

STROUP, W. W. **Generalized Linear Mixed Models: Modern Concepts, Methods and Applications**. 2. ed. Boca Raton: CRC Press, 2018.

Sá, L. R. M. *et al.* Metodologia reml/blup para predição de valor genético em experimentos com correção espacial. **Revista Brasileira de Ciências Agrárias**, v. 18, n. 1, p. e7195, 2023.

TADESE, D.; PHO, H.-P. Spatial model selection and design evaluation in the ethiopian sorghum breeding program. **Agronomy Journal**, p. 2888–2899, 2023.

TANG, W.; ZHANG, L.; BANERJEE, S. On identifiability and consistency of the nugget in gaussian spatial process models. **Journal of the Royal Statistical Society: Series B**, v. 83, n. 5, p. 1044–1070, 2021.

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

WEST, B. T.; WELCH, K. B.; GALECKI, A. T. **Linear Mixed Models: A Practical Guide using Statistical Software**. 3rd. ed. Boca Raton: CRC Press, 2022.

WILKS, S. S. The large-sample distribution of the likelihood ratio for testing composite hypotheses. **Annals of Mathematical Statistics**, v. 9, n. 1, p. 60–62, 1938. Disponível em: <https://doi.org/10.1214/aoms/1177732360>.

WRIGHT, D. **agridat: Agricultural datasets**. [S.l.], 2024. R package version 1.25.0. Disponível em: <https://cran.r-project.org/web/packages/agridat/index.html>.

YWATA, E. d. O.; ALBUQUERQUE, P. H. M. d. S. Testes de hipóteses em modelos lineares mistos: uma abordagem com o proc mixed do sas. **Revista de Economia e Sociologia Rural**, v. 49, n. 2, p. 377–396, 2011.

ZUUR, A. F. *et al.* **Mixed Effects Models and Extensions in Ecology with R**. New York: Springer Science & Business Media, 2009.

APÊNDICES

APÊNDICE A – FUNÇÃO DOUBLEAR1

```
library(MASS)

doubleAR1 <- function(formula, data, start, test_treatment = TRUE,
                      alpha = 0.05, language = "pt") {

  if (!language %in% c("pt", "en")) {
    stop("Language must be 'pt' for Portuguese or 'en' for English")
  }

  if (!all(c("row", "col") %in% names(data))) {
    stop(ifelse(language == "pt",
                "Dados devem conter colunas 'row' e 'col'
                para coordenadas espaciais",
                "Data must contain 'row' and 'col' columns
                for spatial coordinates"))
  }

  if (length(start) != 4 || !all(names(start) %in% c("sigma2",
"tau2", "rho1", "rho2")))) {
    stop(ifelse(language == "pt",
                "start deve ser um vetor nomeado: c(sigma2,
                tau2, rho1, rho2)",
                "start must be a named vector: c(sigma2,
                tau2, rho1, rho2)"))
  }

  texts <- list(
    pt = list(
```

```

warning_rank = "Matriz de design com deficiência de rank.",
  warning_singular = "Usando pseudoinversa par matriz de
covariância devido a singularidade",
  convergence_yes = "SIM",
  convergence_no = "NÃO",
  parameters = c("sigma2", "tau2", "rho_linha", "rho_coluna"),
  descriptions = c(
    "Variância da componente espacial",
    "Variância do efeito pepita",
    "Coeficiente de autocorrelação das linhas",
    "Coeficiente de autocorrelação das colunas"
  ),
  criteria_names = c("AIC (REML)", "BIC (REML)",
    "Log-Verossimilhança REML"),
  significance = c("***", "**", "*", ".", "n.s."),
  wald_title = "TESTE DE WALD MULTIVARIADO - EFEITO DO TRATAMENTO",
  wald_stat = "Estatística Wald",
  wald_df = "Graus de liberdade",
  wald_pval = "P-valor",
  wald_sig = "Significância",
  lrt_title = "TESTE DE RAZÃO DE VEROSSIMILHANÇA (LRT) -
EFEITO DO TRATAMENTO",
  lrt_stat = "Estatística LRT",
  lrt_df = "Graus de liberdade",
  lrt_pval = "P-valor",
  lrt_sig = "Significância",
  pairwise_title = "COMPARAÇÕES PAREADAS DE TRATAMENTOS (TESTE WALD)",
  pairwise_subtitle = "Testes Wald para diferenças
entre pares de tratamentos",
  pairwise_cols = c("Tratamento A", "Tratamento B", "Diferença",
    "EP", "Estatística Wald", "P-valor", "Significância"),
  pairwise_none =

```

```

    "Nenhuma comparação pareada disponível
(sem efeito de tratamento)",
    sections = c(
        "PARÂMETROS DE COVARIÂNCIA",
        "COEFICIENTES DOS EFEITOS FIXOS",
        "COMPARAÇÕES PAREADAS DE TRATAMENTOS",
        "CRITÉRIOS DE INFORMAÇÃO",
        "DIAGNÓSTICOS"
    ),
    diagnostics = c("Convergência", "Número de observações"),
    message_ok = "OK",
    message_collinear = "Algumas colunas removidas
devido a colinearidade"
),
en = list(
    warning_rank = "Design matrix is rank deficient.
Using QR decomposition to select independent columns.",
    warning_singular = "Using pseudoinverse for covariance matrix
due to singularity",
    convergence_yes = "YES",
    convergence_no = "NO",
    parameters = c("sigma2", "tau2", "rho_row", "rho_col"),
    descriptions = c(
        "Spatial component variance",
        "Nugget effect variance",
        "Row autocorrelation coefficient",
        "Column autocorrelation coefficient"
    ),
    criteria_names = c("AIC (REML)", "BIC (REML)",
        "Log-Likelihood REML"),
    significance = c("***", "**", "*", ".", "n.s."),
    wald_title = "MULTIVARIATE WALD TEST - TREATMENT EFFECT",

```

```

wald_stat = "Wald statistic",
wald_df = "Degrees of freedom",
wald_pval = "P-value",
wald_sig = "Significance",
lrt_title = "LIKELIHOOD RATIO TEST (LRT) - TREATMENT EFFECT",
lrt_stat = "LRT statistic",
lrt_df = "Degrees of freedom",
lrt_pval = "P-value",
lrt_sig = "Significance",
pairwise_title = "PAIRWISE TREATMENT COMPARISONS (WALD TEST)",
pairwise_subtitle =
"Wald tests for differences between treatment pairs",
pairwise_cols = c("Treatment A", "Treatment B",
"Difference", "SE", "Wald Statistic",
"P-value", "Significance"),
pairwise_none =
"No pairwise comparisons available (no treatment effect)",
sections = c(
"COVARIANCE PARAMETERS",
"FIXED EFFECTS COEFFICIENTS",
"PAIRWISE TREATMENT COMPARISONS",
"INFORMATION CRITERIA",
"DIAGNÓSTICOS"
),
diagnostics = c("Convergence", "Number of observations"),
message_ok = "OK",
message_collinear = "Some columns removed due to collinearity"
)
)

txt <- texts[[language]]

```

```

mf <- model.frame(formula, data = data)
X_full <- model.matrix(formula, data = mf)
y <- model.response(mf, "numeric")

treatment_removed <- FALSE
reduced_loglik <- NULL
treatment_indices <- NULL
treatment_vars <- NULL
treatment_factor_name <- NULL
treatment_levels <- NULL

treatment_terms <- attr(terms(formula), "term.labels")
has_treatment <- test_treatment && any(grepl("^trat",
treatment_terms, ignore.case = TRUE))

if (has_treatment && test_treatment) {
  treatment_vars <- grep("^trat", treatment_terms,
                        ignore.case = TRUE, value = TRUE)

  if (length(treatment_vars) > 0) {
    treatment_factor_name <- treatment_vars[1]

    if (treatment_factor_name %in% names(data)) {
      if (is.factor(data[[treatment_factor_name]])) {
        treatment_levels <- levels(data[
          treatment_factor_name])
      } else {
        treatment_levels <- unique(as.character(
          data[[treatment_factor_name]]))
      }
    }
  }
}

```

```

response_var <- all.vars(formula)[1]
reduced_terms <- setdiff(treatment_terms, treatment_vars)

if (length(reduced_terms) > 0) {
  reduced_formula <- as.formula(paste(response_var, "~",
                                     paste(reduced_terms, collapse = " + ")))
} else {
  reduced_formula <- as.formula(paste(response_var, "~ 1"))
}

treatment_removed <- TRUE

reduced_model <- tryCatch({
  doubleAR1(reduced_formula, data, start, test_treatment = FALSE,
            alpha = alpha, language = language)
}, error = function(e) {
  warning(ifelse(language == "pt",
                 "Não foi possível ajustar
                 o modelo reduzido para LRT",
                 "Could not fit reduced model for LRT"))

  NULL
})

if (!is.null(reduced_model)) {
  reduced_loglik <- reduced_model$log_verossimilhanca
} } }

X_REML <- X_full

if (has_treatment && test_treatment && length(treatment_vars) > 0) {
  treatment_cols <- grep(paste(treatment_vars, collapse = "|"),
                        colnames(X_REML))

```

```

    if (length(treatment_cols) > 0) {
      treatment_indices <- treatment_cols
    }
  }

  if (ncol(X_REML) > 0) {
    X_rank <- qr(X_REML)$rank
    if (X_rank < ncol(X_REML)) {
      warning(txt$warning_rank)
      qr_decomp <- qr(X_REML)
      pivot <- qr_decomp$pivot[1:qr_decomp$rank]
      X_REML <- X_REML[, pivot, drop = FALSE]

      if (!is.null(treatment_indices)) {
        treatment_indices <- which(pivot %in% treatment_indices)
      }

      colnames_original <- colnames(X_full)[pivot]
      colnames(X_REML) <- colnames_original
      message_collinear <- TRUE
    }
  }

  N <- nrow(data)
  P <- ncol(X_REML)
  r <- max(data$row)
  c <- max(data$col)

  construir_V <- function(sigma2, tau2, rho1, rho2, r, c) {
    Rlinha <- outer(1:r, 1:r, function(i, j) rho1^abs(i - j))
    Rcoluna <- outer(1:c, 1:c, function(i, j) rho2^abs(i - j))
    V <- sigma2 * kronecker(Rlinha, Rcoluna) + tau2 * diag(N)
  }

```

```

    return(V)
}

neg_log_likelihoood_reml <- function(params, X, y, N, r, c, txt) {
  sigma2 <- params[1]
  tau2 <- params[2]
  rho1 <- params[3]
  rho2 <- params[4]

  P <- ncol(X)

  if (sigma2 < 1e-6 || tau2 < 1e-6 || abs(rho1) >= 1 ||
      abs(rho2) >= 1)
  {
    return(1e20)}

  V <- construire_V(sigma2, tau2, rho1, rho2, r, c)

  tryCatch({
    V_inv <- tryCatch(solve(V), error = function(e) {
      warning(txt$warning_singular)
      MASS::ginv(V)
    })
    log_det_V <- as.numeric(determinant(
      V, logarithm = TRUE)$modulus)

    X_Vinv_X <- t(X) %*% V_inv %*% X

    X_Vinv_X_inv <- tryCatch(solve(X_Vinv_X), error = function(e) {
      warning(txt$warning_singular)
      MASS::ginv(X_Vinv_X)
    })
  })
}

```

```

log_det_XVinvX <- as.numeric(
determinant(X_Vinv_X, logarithm = TRUE)$modulus)

B_component <- X %*% X_Vinv_X_inv %*% t(X) %*% V_inv

P_matrix <- V_inv - V_inv %*% B_component

y_P_y <- as.numeric(t(y) %*% P_matrix %*% y)

l_R <- -0.5 * (
  (N - P) * log(2 * pi) +
  log_det_V +
  log_det_XVinvX +
  y_P_y
)

return(-l_R)

}, error = function(e) {
  return(1e20)
})
}

resultado_bfgs <- optim(
  par = start,
  fn = neg_log_likelihood_reml,
  method = "L-BFGS-B",
  lower = c(sigma2 = 1e-6, tau2 = 1e-6, rho1 = -0.999,
rho2 = -0.999),
  upper = c(sigma2 = Inf, tau2 = Inf, rho1 = 0.999, rho2 = 0.999),
  X = X_REML, y = y, N = N, r = r, c = c, txt = txt

```

```

)

sigma2_est <- resultado_bfgs$par[1]
tau2_est <- resultado_bfgs$par[2]
rho1_est <- resultado_bfgs$par[3]
rho2_est <- resultado_bfgs$par[4]

V_est <- construir_V(sigma2_est, tau2_est, rho1_est,
rho2_est, r, c)
V_inv_est <- tryCatch(solve(V_est),
error = function(e) MASS::ginv(V_est))

X_Vinv_X_est <- t(X_REML) %*% V_inv_est %*% X_REML

beta_hat <- tryCatch({
  solve(X_Vinv_X_est) %*% t(X_REML) %*% V_inv_est %*% y
}, error = function(e) {
  MASS::ginv(X_Vinv_X_est) %*% t(X_REML) %*% V_inv_est %*% y
})

cov_beta <- tryCatch({
  solve(X_Vinv_X_est)
}, error = function(e) {
  warning(txt$warning_singular)
  MASS::ginv(X_Vinv_X_est)
})

se_beta <- sqrt(diag(cov_beta))

estatistica_wald <- rep(NA, length(beta_hat))
p_valor_wald <- rep(NA, length(beta_hat))

```

```

valid_se <- se_beta > 1e-10
if (any(valid_se)) {
  estatistica_wald[valid_se] <- (
    beta_hat[valid_se] / se_beta[valid_se])^2
  p_valor_wald[valid_se] <- 1 - pchisq(estatistica_wald[valid_se],
    df = 1)
}

wald_result <- NULL
if (!is.null(treatment_indices) && length(treatment_indices) > 0) {
  s <- length(treatment_indices)
  h <- length(beta_hat)

  L <- matrix(0, nrow = s, ncol = h)
  for (i in 1:s) {
    L[i, treatment_indices[i]] <- 1
  }

  c_vector <- rep(0, s)

  L_beta <- L %*% beta_hat
  L_cov_Lt <- L %*% cov_beta %*% t(L)

  tryCatch({
    wald_statistic <- t(L_beta - c_vector) %*%
      solve(L_cov_Lt) %*% (L_beta - c_vector)
    wald_statistic <- as.numeric(wald_statistic)

    p_valor_wald_multi <- pchisq(wald_statistic,
      df = s, lower.tail = FALSE)

    wald_result <- list(

```

```

    Wald_statistic = wald_statistic,
    Wald_df = s,
    Wald_p_value = p_valor_wald_multi,
    treatment_indices = treatment_indices,
    treatment_terms = colnames(X_REML)[treatment_indices]
  )
}, error = function(e) {
  warning(ifelse(language == "pt",
    "Não foi possível calcular o teste de Wald multivariado",
    "Could not compute multivariate Wald test"))
})
}

pairwise_comparisons <- NULL
if (has_treatment && !is.null(treatment_levels) &&
  length(treatment_levels) > 1) {
  k <- length(treatment_levels)
  treatment_coef_map <- list()
  treatment_coef_map[[treatment_levels[1]]] <- 0

  for (i in 2:k) {
    trt_name <- treatment_levels[i]
    coef_pattern <- paste0("^", treatment_factor_name,
      trt_name, "$")
    coef_idx <- grep(coef_pattern, colnames(X_REML))
    if (length(coef_idx) == 1) {
      treatment_coef_map[[trt_name]] <- coef_idx
    } else {
      coef_pattern2 <- paste0("^", trt_name, "$")
      coef_idx <- grep(coef_pattern2, colnames(X_REML))
      if (length(coef_idx) == 1) {
        treatment_coef_map[[trt_name]] <- coef_idx
      }
    }
  }
}

```

```

    } else {
      treatment_coef_map[[trt_name]] <- NA
    }
  }
}

treatment_A <- character(0)
treatment_B <- character(0)
differences <- numeric(0)
standard_errors <- numeric(0)
wald_stats <- numeric(0)
p_values <- numeric(0)

for (i in 1:(k-1)) {
  for (j in (i+1):k) {
    trt_A <- treatment_levels[i]
    trt_B <- treatment_levels[j]

    idx_A <- treatment_coef_map[[trt_A]]
    idx_B <- treatment_coef_map[[trt_B]]

    if (idx_A == 0) {
      mean_A <- beta_hat[1]
      if (!is.na(idx_B) && idx_B > 0) {
        mean_B <- beta_hat[1] + beta_hat[idx_B]
        diff_est <- mean_A - mean_B
        var_diff <- cov_beta[idx_B, idx_B]
      } else {
        diff_est <- NA
        var_diff <- NA
      }
    } else if (idx_B == 0) {

```

```

mean_B <- beta_hat[1]
if (!is.na(idx_A) && idx_A > 0) {
  mean_A <- beta_hat[1] + beta_hat[idx_A]
  diff_est <- mean_A - mean_B
  var_diff <- cov_beta[idx_A, idx_A]
} else {
  diff_est <- NA
  var_diff <- NA
}
} else if (
!is.na(idx_A) && !is.na(idx_B) && idx_A > 0 && idx_B > 0)
{
  mean_A <- beta_hat[1] + beta_hat[idx_A]
  mean_B <- beta_hat[1] + beta_hat[idx_B]
  diff_est <- mean_A - mean_B
  var_diff <- cov_beta[idx_A, idx_A] +
    cov_beta[idx_B, idx_B] - 2 * cov_beta[idx_A, idx_B]
} else {
  diff_est <- NA
  var_diff <- NA
}

if (!is.na(var_diff) && var_diff > 0) {
  se_diff <- sqrt(var_diff)
  if (se_diff > 1e-10 && !is.na(diff_est)) {
    wald_stat_diff <- (diff_est / se_diff)^2
    p_val_diff <- 1 - pchisq(wald_stat_diff, df = 1)
  } else {
    wald_stat_diff <- NA
    p_val_diff <- NA
  }
} else {

```

```

    se_diff <- NA
    wald_stat_diff <- NA
    p_val_diff <- NA
  }

  treatment_A <- c(treatment_A, trt_A)
  treatment_B <- c(treatment_B, trt_B)
  differences <- c(differences,
    ifelse(is.na(diff_est), NA, diff_est))
  standard_errors <- c(standard_errors, se_diff)
  wald_stats <- c(wald_stats, wald_stat_diff)
  p_values <- c(p_values, p_val_diff)
}

}

significance <- character(length(p_values))
for (i in 1:length(p_values)) {
  p_val <- p_values[i]
  if (is.na(p_val)) {
    significance[i] <- "NA"
  } else if (p_val < 0.001) {
    significance[i] <- "****"
  } else if (p_val < 0.01) {
    significance[i] <- "***"
  } else if (p_val < 0.05) {
    significance[i] <- "**"
  } else if (p_val < 0.1) {
    significance[i] <- "."
  } else {
    significance[i] <- "n.s."
  }
}

```

```

pairwise_comparisons <- data.frame(
  Treatment_A = treatment_A,
  Treatment_B = treatment_B,
  Difference = differences,
  SE = standard_errors,
  Wald_Statistic = wald_stats,
  P_Value = p_values,
  Significance = significance,
  stringsAsFactors = FALSE
)

if (nrow(pairwise_comparisons) > 0) {
pairwise_comparisons <- pairwise_comparisons[order(
pairwise_comparisons$P_Value)
, ]

} }

log_verossimilhanca_reml <- -resultado_bfgs$value

num_param_variance <- 4
n_eff <- N - P
AIC_reml <- 2 * num_param_variance - 2 * log_verossimilhanca_reml
BIC_reml <- log(N) * num_param_variance -
2 * log_verossimilhanca_reml

lrt_result <- NULL
if (!is.null(reduced_loglik) && treatment_removed) {
  lrt_statistic <- 2 *
    (log_verossimilhanca_reml - reduced_loglik)
  lrt_statistic <- max(0, lrt_statistic)

  treatment_factor <- data[[treatment_factor_name]]

```

```

if (is.factor(treatment_factor)) {
  df_lrt <- nlevels(treatment_factor) - 1
} else {
  df_lrt <- 1
}

p_valor_lrt <- pchisq(lrt_statistic, df = df_lrt,
lower.tail = FALSE)

lrt_result <- list(
  LRT_statistic = lrt_statistic,
  LRT_df = df_lrt,
  LRT_p_value = p_valor_lrt,
  treatment_removed = treatment_removed
)
}

tabela_parametros <- data.frame(
  Parameter = txt$parameters,
  Estimate = c(sigma2_est, tau2_est, rho1_est, rho2_est),
  Description = txt$descriptions
)

nomes_coef <- colnames(X_REML)

tabela_coeficientes <- data.frame(
  Coefficient = nomes_coef,
  Estimate = as.numeric(beta_hat),
  Std_Error = se_beta,
  Wald_Statistic = as.numeric(estatistica_wald),
  P_Value_Wald = as.numeric(p_valor_wald)
)

```

```

tabela_criterios <- data.frame(
  Criterion = txt$criteria_names,
  Value = c(AIC_reml, BIC_reml, log_verossimilhanca_reml)
)

residuos <- y - X_REML %*% beta_hat

resultado_completo <- list(
  treatment_wald = wald_result,
  treatment_lrt = lrt_result,
  pairwise_comparisons = pairwise_comparisons,
  covariance_parameters = tabela_parametros,
  fixed_effects = tabela_coeficientes,
  information_criteria = tabela_criterios,
  wald_tests = list(
    statistics = estatistica_wald,
    p_values = p_valor_wald
  ),
  convergence = resultado_bfgs$convergence,
  objective_value = resultado_bfgs$value,
  residuals = as.numeric(residuos),
  V_matrix = V_est,
  log_verossimilhanca = log_verossimilhanca_reml,
  language = language,
  texts = txt,
  X_matrix_used = X_REML,
  alpha = alpha,
  treatment_levels = treatment_levels,
  message = if(exists("message_collinear")) txt$message_collinear
  else txt$message_ok
)

```

```

class(resultado_completo) <- "doubleAR1"
return(resultado_completo)
}

print.doubleAR1 <- function(x, ...) {
  txt <- x$texts

  if (x$language == "pt") {
    cat("===      MODELO AR(1) × AR(1) -
    RESULTADOS COM TESTE DE WALD ===\n\n")
  } else {
    cat("=== AR(1) × AR(1) MODEL -
    RESULTS WITH WALD TEST ===\n\n")
  }

  if (!is.null(x$message) && x$message != txt$message_ok) {
    cat("AVISO:", x$message, "\n\n")
  }

  if (!is.null(x$treatment_wald)) {
    cat(txt$wald_title, ":\n")
    cat("  ", txt$wald_stat, ":", x$treatment_wald$Wald_statistic,
    "\n")
    cat("  ", txt$wald_df, ":", x$treatment_wald$Wald_df, "\n")
    cat("  ", txt$wald_pval, ":",
      format(x$treatment_wald$Wald_p_value,
      scientific = FALSE), "\n")
    cat("  Termos testados:",
      paste(x$treatment_wald$treatment_terms,
      collapse = ", "), "\n")
    cat("\n\n")
  }
}

```

```

if (!is.null(x$treatment_lrt)) {
  cat(txt$lrt_title, ":\n")
  cat("  ", txt$lrt_stat, ":", x$treatment_lrt$LRT_statistic,
    "\n")
  cat("  ", txt$lrt_df, ":", x$treatment_lrt$LRT_df, "\n")
  cat("  ", txt$lrt_pval, ":",
    format(x$treatment_lrt$LRT_p_value,
      scientific = FALSE), "\n")
  cat("\n")
}

for (i in 1:5) {
  cat(i, ". ", txt$sections[i], ":\n", sep = "")
  switch(i, print(x$covariance_parameters),
    {if (!is.null(x$fixed_effects) && nrow(x$fixed_effects) > 0) {
      fe_table <- x$fixed_effects
      if ("Coefficient" %in% colnames(fe_table)) {
        rownames(fe_table) <- fe_table$Coefficient
        fe_table <- fe_table[, -1, drop = FALSE]
      }
      colnames(fe_table) <- c("Estimativa", "EP",
        "Estatística Wald", "P-valor")
      if (nrow(fe_table) > 20) {
        cat("  (Mostrando apenas os primeiros 20 coeficientes)\n")
        print(head(fe_table, 20))
        cat("    ... [" , nrow(fe_table) - 20,
          " coeficientes omitidos]\n")
      } else {
        print(fe_table)
      }
    } else {
  }
}

```

```

        cat(" Nenhum coeficiente disponível\n")
      } },
    {
      if (!is.null(x$pairwise_comparisons) &&
          nrow(x$pairwise_comparisons) > 0) {
        cat(" ", txt$pairwise_subtitle, "\n")
        cat("  $\alpha$  =", x$alpha, "\n")
        cat(" Número total de tratamentos:", length(x$treatment_levels), "\n")
        cat(" Número total de combinações dois a dois:",
              nrow(x$pairwise_comparisons), "\n\n")

        pairwise_display <- data.frame(
          A = x$pairwise_comparisons$Treatment_A,
          B = x$pairwise_comparisons$Treatment_B,
          Dif = format(round(x$pairwise_comparisons$Difference, 4),
                        nsmall = 4),
          EP = format(round(x$pairwise_comparisons$SE, 4), nsmall = 4),
          Wald = format(round(x$pairwise_comparisons$Wald_Statistic, 3),
                         nsmall = 3),
          P = format(x$pairwise_comparisons$P_Value, scientific = TRUE,
                      digits = 3),
          Sig = x$pairwise_comparisons$Significance
        )

        colnames(pairwise_display) <- txt$pairwise_cols
        print(pairwise_display, row.names = FALSE)

        cat("\n RESUMO:\n")
        cat(" Número total de comparações:",
              nrow(x$pairwise_comparisons), "\n")
        sig_count <- sum(
          x$pairwise_comparisons$P_Value < x$alpha, na.rm = TRUE)
        cat(" Comparações significativas ( $\alpha$  =", x$alpha, "):",

```

```

sig_count, "\n")
cat(" Percentual significativo:",
    round(sig_count/nrow(x$pairwise_comparisons)*100, 1), "%\n")
    } else {
        cat(" ", txt$pairwise_none, "\n")
    }},
    print(x$information_criteria),
{
    cat(txt$diagnostics[1], ":",
        ifelse(x$convergence == 0,
            txt$convergence_yes, txt$convergence_no), "\n")
        cat(txt$diagnostics[2], ":", length(x$residuals), "\n")
    } )
    cat("\n")
}}

summary.doubleAR1 <- function(object, ...) {
    print.doubleAR1(object, ...)
}

pairwise_comparisons <- function(x, ...) {
    UseMethod("pairwise_comparisons")
}

pairwise_comparisons.doubleAR1 <- function(x, ...) {
    if (is.null(x$pairwise_comparisons)) {
        warning("No pairwise comparisons available in this model")
        return(NULL)
    }
    return(x$pairwise_comparisons)
}

```

```

significant_comparisons <- function(x, alpha = NULL, ...) {
  UseMethod("significant_comparisons")
}

significant_comparisons.doubleAR1 <- function(x, alpha = NULL, ...) {
  if (is.null(x$pairwise_comparisons)) {
    warning("No pairwise comparisons available in this model")
    return(NULL)
  }
  if (is.null(alpha)) {
    alpha <- x$alpha
  }
  sig_comparisons <- x$pairwise_comparisons[
    x$pairwise_comparisons$P_Value < alpha &
    !is.na(x$pairwise_comparisons$P_Value),
  ]
  return(sig_comparisons)
}

```

APÊNDICE B – SCRIPT EXEMPLO DIC

```

set.seed(123)
#SIMULAÇÃO DOS DADOS EM ESTRUTURA DE GRID (LATTICE)
n_linhas <- 4
n_colunas <- 8
n_total <- n_linhas * n_colunas
dados_sim <- expand.grid(
  Linha = 1:n_linhas,
  Coluna = 1:n_colunas
)
dados_sim$Tratamento <- factor(sample(
  rep(c("T1", "T2", "T3", "T4"), times = 8),
  size = n_total,
  replace = FALSE
))
dados_sim <- dados_sim[order(dados_sim$Linha, dados_sim$Coluna), ]
efeito_tratamento <- 0.01 * as.numeric(dados_sim$Tratamento)
phi_row <- 0.98
phi_col <- 0.98
sigma_erro <- 20
erro_espacial <- numeric(n_total)
for (i in 1:n_linhas) {
  for (j in 1:n_colunas) {
    idx <- which(dados_sim$Linha == i & dados_sim$Coluna == j)
    erro_espacial[idx] <- (i * phi_row + j * phi_col) * sigma_erro
  }
}
ruído_aleatorio <- rnorm(n_total, mean = 0, sd = 0.5)
dados_sim$Resposta <- (
  100 + efeito_tratamento + erro_espacial + ruído_aleatorio)
modelo_simples <- aov(Resposta ~ Tratamento, data = dados_sim)

```

```
anova(modelo_simples)
BIC(modelo_simples)
AIC(modelo_simples)

var_total <- var(dados_sim$Resposta)

start_values <- c(
  sigma2 = var_total*0.7,
  tau2 = var_total*0.3,
  rho1 = 0.95,
  rho2 = 0.95
)
names(dados_sim)[names(dados_sim) == "Linha"] <- "row"
names(dados_sim)[names(dados_sim) == "Coluna"] <- "col"
resultado_dic <- doubleAR1(
  Resposta ~ Tratamento,
  data = dados_sim,
  start = start_values,
  test_treatment = TRUE,
  language = "pt"
)
print(resultado_dic)
```

APÊNDICE C – SCRIPT EXEMPLO DBC

```

library(nlme)
library(emmeans)
library(agricolae)

data(eden.potato)
dat_dbc <- subset(eden.potato, year == '1926b')
dat_dbc <- droplevels(dat_dbc)
names(dat_dbc)[names(dat_dbc) == "trt"] <- "trat"
var_total <- var(dat_dbc$yield)
start_values <- c(
  sigma2 = var_total * 0.7,
  tau2 = var_total * 0.3,
  rho1 = 0.2,
  rho2 = 0.2
)

modelo_completo <- doubleAR1(
  formula = yield ~ trat + block,
  data = dat_dbc,
  start = start_values,
  test_treatment = TRUE,
  language = "pt"
)

# Ver tudo
modelo_completo

# Ver apenas o teste de Wald multivariado
modelo_completo$treatment_wald

```

```
# Ver apenas o teste LRT
```

```
modelo_completo$treatment_lrt
```

```
# Ver apenas os parâmetros de covariância
```

```
modelo_completo$covariance_parameters
```

```
# Ver apenas a comparação dos tratamentos dois a dois
```

```
modelo_completo$pairwise_comparisons
```

```
# Ver apenas os critérios de informação
```

```
modelo_completo$information_criteria
```