



Haward Antunny da Silva Américo

**AVALIAÇÃO DO AJUSTE DE MODELOS NÃO
LINEARES MISTOS À ANÁLISE DE
DEGRADABILIDADE RUMINAL *IN VITRO***

Maringá - Paraná
2026

Haward Antunny da Silva Américo

**AVALIAÇÃO DO AJUSTE DE MODELOS NÃO LINEARES
MISTOS À ANÁLISE DE DEGRADABILIDADE RUMINAL
*IN VITRO***

Dissertação apresentada ao Programa de Pós-graduação em Bioestatística do Centro de Ciências Exatas da Universidade Estadual de Maringá como requisito parcial para obtenção do título de mestre em Bioestatística.

Orientador: Prof. Dr. Vanderly Janeiro, Universidade Estadual de Maringá.

Membro do PBE: Prof. Dr. Marcos Vinicius de Oliveira Peres, Universidade Estadual de Maringá.

Membro Externo: Prof. Dr. Ricardo Souza Vasconcellos, Universidade Estadual de Maringá.

Universidade Estadual de Maringá - UEM

Departamento de Estatística - DES

Programa de Pós-Graduação em Bioestatística

Maringá - Paraná

2026

Dados Internacionais de Catalogação-na-Publicação (CIP)
(Biblioteca Central - UEM, Maringá - PR, Brasil)

A512a

Américo, Haward Antunny da Silva

Avaliação do ajuste de modelos não lineares mistos à análise de degradabilidade ruminal *in vitro* / Haward Antunny da Silva Américo. -- Maringá, PR, 2026.
100 f. : il. color., figs., tabs.

Orientador: Prof. Dr. Vanderly Janeiro.

Dissertação (mestrado) - Universidade Estadual de Maringá, Centro de Ciências Exatas, Departamento de Estatística, Programa de Pós-Graduação em Bioestatística, 2026.

1. Modelos não lineares mistos. 2. Correlação temporal. 3. Produção de gás *in vitro*. 4. Degradabilidade efetiva. 5. Estruturas de variância covariância. I. Janeiro, Vanderly, orient. II. Universidade Estadual de Maringá. Centro de Ciências Exatas. Departamento de Estatística. Programa de Pós-Graduação em Bioestatística. III. Título.

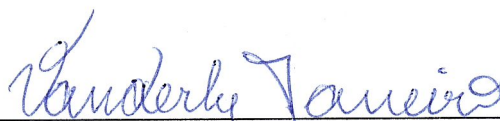
CDD 23.ed. 519.5

HAWARD ANTUNNY DA SILVA AMÉRICO

**Avaliação do Ajuste de Modelos Não Lineares Mistos à Análise de
Degradabilidade Ruminal *in vitro***

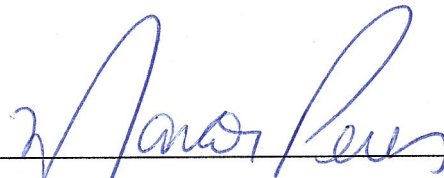
Dissertação apresentada ao Programa de Pós-Graduação em Bioestatística do Centro de Ciências Exatas da Universidade Estadual de Maringá, como requisito parcial para a obtenção do título de Mestre em Bioestatística.

BANCA EXAMINADORA



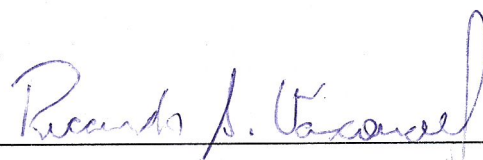
Prof. Dr. Vanderly Janeiro (Presidente)

Universidade Estadual de Maringá – PBE/UEM



Prof. Dr. Marcos Vinícius de Oliveira Peres

Universidade Estadual de Maringá – PBE/UEM



Prof. Dr. Ricardo Souza Vasconcellos

Universidade Estadual de Maringá – PPZ/UEM

Maringá, 24 de dezembro de 2026.

AGRADECIMENTOS

À minha mãe, cuja presença firme e amorosa acompanhou cada etapa da minha vida e formação. Mulher forte, batalhadora e resiliente, que sempre foi meu exemplo diário de coragem e de dedicação. Seu apoio incondicional, suas palavras de ânimo e sua fé na minha capacidade sustentaram-me nos momentos de incerteza, nas mudanças inesperadas e nos inúmeros desafios enfrentados ao longo destes anos. É a ela que devo grande parte do que sou e do que conquistei.

À Prof^a. Dr^a. Lucimary Afonso, minha orientadora na graduação, que acreditou no meu potencial muito antes de eu ingressar no mestrado. Sua orientação paciente e incentivo constante despertaram em mim o desejo de seguir na pós-graduação e aprofundar-me na estatística aplicada e na pesquisa científica.

Ao meu orientador de mestrado, Prof. Dr. Vanderly Janeiro, que me acompanha desde a graduação e que aceitou caminhar comigo neste novo desafio. Sua disponibilidade, seu rigor acadêmico, sua generosidade ao ensinar e sua confiança no meu trabalho foram fundamentais para o meu crescimento. Sou profundamente grato por cada sugestão, cada artigo recomendado e cada momento dedicado a esclarecer dúvidas e ampliar minha formação. Seu exemplo de profissionalismo e ética é uma inspiração contínua.

Aos pesquisadores com quem tive a oportunidade de trabalhar ao longo do programa. Cada projeto desenvolvido trouxe novos aprendizados, desafios e descobertas que contribuíram significativamente para minha formação como bioestatístico. Estendo também minha gratidão aos amigos e colegas da pós-graduação, que compartilharam comigo tantas horas de estudo na “salinha”, trocas de conhecimento, incentivos mútuos e momentos essenciais de parceria acadêmica e pessoal.

À CAPES, pelo fomento da bolsa de estudos, que tornou possível minha dedicação ao mestrado. O apoio financeiro proporcionado pela instituição foi essencial para que eu pudesse desenvolver este trabalho com a profundidade e o compromisso que ele exige.

Aos membros desta banca avaliadora, Prof. Dr. Marcos Vinicius de Oliveira Peres e Prof. Dr. Ricardo Souza Vasconcellos, pela gentileza em aceitar o convite para avaliar este trabalho. Suas leituras atentas, críticas construtivas e contribuições certamente enriquecerão ainda mais a qualidade desta pesquisa.

Por fim, agradeço a todos que, direta ou indiretamente, fizeram parte desta trajetória. Cada palavra de incentivo, cada gesto de apoio e cada experiência compartilhada contribuíram para que este mestrado se tornasse possível.

RESUMO

A digestibilidade da fração fibrosa em ruminantes influencia diretamente o desempenho produtivo e a eficiência alimentar. A técnica de produção de gás *in vitro* é amplamente utilizada para avaliar a fermentabilidade ruminal de alimentos, gerando dados longitudinais com medidas repetidas, dependência temporal e variabilidade entre unidades experimentais. Nesse contexto, a utilização de modelos estatísticos que incorporem simultaneamente efeitos aleatórios e estruturas adequadas de covariância residual é fundamental para garantir inferências estatisticamente consistentes e biologicamente interpretáveis. O objetivo deste estudo foi ajustar, comparar e avaliar modelos não lineares aplicados à produção cumulativa de gás, considerando as funções de Ørskov e McDonald, Gompertz e Logística, sob abordagens de efeitos fixos e modelos não lineares mistos. As análises foram realizadas no ambiente R, por meio do pacote *nlme*, com investigação sistemática das estruturas de efeitos aleatórios e das matrizes de variância covariância dos parâmetros (**D**) e dos erros (**R**). A seleção dos modelos baseou-se em critérios de informação (AIC e BIC), testes da razão de verossimilhança, diagnósticos residuais, análise da estrutura de autocorrelação e avaliação do desempenho preditivo, considerando tanto o conjunto de dados utilizado no ajuste quanto um conjunto independente de validação. Foram empregadas métricas como erro médio, raiz do erro quadrático médio, correlação de Pearson e coeficiente de concordância de Lin. Análises preliminares conduzidas sem a especificação adequada da estrutura de covariância residual evidenciaram instabilidade nas estimativas, incluindo valores biologicamente inconsistentes para alguns parâmetros. A incorporação de estruturas de correlação residual mais flexíveis, em especial funções contínuas de dependência temporal, resultou em estimativas mais estáveis, coerentes e com melhor desempenho preditivo. Entre os modelos avaliados, o modelo de Ørskov e McDonald com efeitos aleatórios nos parâmetros β_1 , β_2 e β_3 e estrutura adequada para a matriz residual apresentou o melhor desempenho global, considerando simultaneamente critérios de ajuste, estabilidade das estimativas e capacidade preditiva. As estimativas por tratamento permitiram o cálculo da degradabilidade efetiva, evidenciando diferenças associadas principalmente à fração potencialmente degradável. Conclui-se que os modelos não lineares mistos, com especificação adequada das estruturas de variância e correlação, constituem uma abordagem robusta, parcimoniosa e biologicamente consistente para a análise de dados de produção de gás *in vitro*.

Palavras-chave: Modelos não lineares mistos, Produção de gás *in vitro*, Estruturas de variância covariância, Correlação temporal, Degradabilidade efetiva.

ABSTRACT

The digestibility of the fibrous fraction in ruminants directly influences productive performance and feed efficiency. The *in vitro* gas production technique is widely used to evaluate the ruminal fermentability of feedstuffs, generating longitudinal data with repeated measurements, temporal dependence, and variability among experimental units. In this context, the use of statistical models that simultaneously incorporate random effects and appropriate residual covariance structures is essential to ensure statistically consistent and biologically interpretable inferences. The objective of this study was to fit, compare, and evaluate nonlinear models applied to cumulative gas production, considering the Ørskov and McDonald, Gompertz, and Logistic functions under both fixed-effects and nonlinear mixed-effects approaches. Analyses were performed in the R environment using the nlme package, with a systematic investigation of random-effects structures and the variance covariance matrices of the parameters (**D**) and the residual errors (**R**). Model selection was based on information criteria (AIC and BIC), likelihood ratio tests, residual diagnostics, autocorrelation structure assessment, and predictive performance evaluation, considering both the training dataset and an independent validation dataset. The metrics included mean error, root mean squared error, Pearson correlation, and Lin's concordance correlation coefficient. Preliminary analyses conducted without appropriate specification of the residual covariance structure revealed instability in parameter estimates, including biologically inconsistent values. The incorporation of more flexible residual correlation structures, particularly continuous-time correlation functions, resulted in more stable, coherent, and predictive models. Among the evaluated models, the Ørskov and McDonald model with random effects on parameters β_1 , β_2 , and β_3 , combined with an appropriate residual covariance structure, showed the best overall performance in terms of goodness-of-fit, parameter stability, and predictive ability. Treatment-level parameter estimates allowed the calculation of effective degradability, highlighting differences mainly associated with the potentially degradable fraction. It is concluded that nonlinear mixed-effects models, with appropriate specification of variance and correlation structures, constitute a robust, parsimonious, and biologically consistent approach for the analysis of *in vitro* gas production data.

Keywords: Nonlinear mixed-effects models, *In vitro* gas production, Variance covariance structures, Temporal correlation, Effective degradability.

LISTA DE ILUSTRAÇÕES

Figura 1 – Animal canulado utilizado para a coleta de fluido ruminal, empregado como inóculo biológico em ensaios de degradabilidade ruminal <i>in vitro</i>	53
Figura 2 – Produção de gás corrigido ao longo do tempo para as amostras.	61
Figura 3 – Curvas de produção de gás excluídas das análises.	62
Figura 4 – Curvas de produção de gás por réplica que permaneceram para ajuste (após exclusão de amostras inconsistentes)	64
Figura 5 – Intervalos de confiança (95%) dos parâmetros estimados para cada amostra, modelo de Ørskov & McDonald (OR). ($A = \beta_1$, $B = \beta_2$, $C = \beta_3$).	67
Figura 6 – Intervalos de confiança (95%) dos parâmetros estimados para cada amostra, modelo de Gompertz (GO). ($A = \beta_1$, $B = \beta_2$, $C = \beta_3$).	68
Figura 7 – Intervalos de confiança (95%) dos parâmetros estimados para cada amostra, modelo Logístico (LO). ($A = \beta_1$, $B = \beta_2$, $C = \beta_3$).	68
Figura 8 – Curvas ajustadas para cada amostra segundo os modelos não lineares fixos de Ørskov, Gompertz e Logístico.	70
Figura 9 – Gráficos de resíduos padronizados em função dos valores ajustados dos MNLM	76
Figura 10 – Autocorrelação dos resíduos dos modelos não lineares mistos por estrutura de tempo	77
Figura 11 – Gráficos HNP dos resíduos dos modelos não lineares mistos ajustados. . . .	80
Figura 12 – Gráficos de resíduos padronizados versus valores ajustados dos modelos com estrutura gaussiana e variânciasheterogêneas.	81
Figura 13 – Gráficos de valores preditos <i>versus</i> observados para os modelos não lineares mistos, considerando predições no nível marginal (level = 0) e condicional (level = 1), avaliados em conjunto de dados de teste.	83
Figura 14 – Valores observados e curvas ajustadas por réplica e tratamento segundo o modelo OR _{7.2.6}	85
Figura 15 – Curvas de produção de gás e pontos de truncamento para os tratamentos avaliados de 1 a 4.	94

Figura 16 – Curvas de produção de gás e pontos de truncamento para os tratamentos avaliados de 5 a 10.	95
Figura 17 – Curvas de produção de gás e pontos de truncamento para os tratamentos avaliados de 11 a 16.	96

LISTA DE TABELAS

Tabela 1 – Métricas utilizadas para avaliação do desempenho preditivo.	52
Tabela 2 – Modelos não lineares mistos com inclusão explícita de termos de efeitos aleatórios nos parâmetros.	57
Tabela 3 – Estatísticas descritivas gerais da produção de gás (G) nos bancos completo e filtrado	64
Tabela 4 – Comparação descritiva dos parâmetros estimados nos modelos não lineares fixos	66
Tabela 5 – Critérios de seleção dos modelos não lineares mistos com diferentes especificações de efeitos aleatórios.	73
Tabela 6 – Comparação das estruturas da matriz de variância-covariância dos efeitos aleatórios ($\mathbf{D}$) para os modelos com efeitos aleatórios em β_1 , β_2 e β_3	75
Tabela 7 – Critérios de seleção dos modelos não lineares mistos com diferentes estruturas da matriz de covariância dos erros ($\mathbf{R}$).	78
Tabela 8 – Métricas preditivas globais dos modelos não lineares mistos ajustados, avaliadas em conjunto de dados de teste.	82
Tabela 9 – Estimativas dos parâmetros do modelo de Ørskov e McDonald ($\text{OR}_{7.2.6}$) e degradabilidade efetiva (DE).	84
Tabela 10 – Componentes da matriz de variância-covariância dos efeitos aleatórios ($\mathbf{D}$) e variância residual para os modelos não lineares mistos, considerando efeitos aleatórios em β_1 , β_2 e β_3	97
Tabela 11 – Componentes estimados associados às diferentes estruturas da matriz de correlação residual ($\mathbf{R}$) nos modelos não lineares mistos.	98
Tabela 12 – Indicadores de instabilidade estrutural no modelo OR_7 com $\mathbf{D}$ geral (pdSymm) e estruturas de variância GAUS + ID.	99
Tabela 13 – Contrastes significativos no teste de Tukey ($p < 0,05$) para o parâmetro β_2 do modelo de Ørskov e McDonald ($\text{OR}_{7.2.6}$).	100

LISTA DE ABREVIATURAS E SIGLAS

ACF	Autocorrelação dos Resíduos Condicionais
AIC	Critério de Informação de Akaike
AICc	Critério de Informação de Akaike Corrigido
ANOVA	Análise de Variância Univariada com Medidas Repetidas
BIC	Critério de Informação Bayesiano
CCC	Coefficiente de Correlação de Concordância
CEUA	Comissão de Ética no Uso de Animais
Corr	Coefficiente de correlação de Pearson
CS	<i>Compound Symmetry</i> (Simetria Composta)
CV	Coefficiente de Variação
DE	Degradabilidade efetiva
DP	Desvio Padrão
EXP	Estrutura de correlação exponencial
ID	Estrutura com variâncias heterogêneas e erros independentes
GAUS	Estrutura de correlação gaussiana
GLS	Mínimos Quadrados Generalizados
GO	Modelo de Gompertz
LO	Modelo Logístico
logLik	Logaritmo da função de verossimilhança
MANOVA	Análise de Variância Multivariada
MAE	Erro Absoluto Médio
MAPE	Erro Percentual Absoluto Médio
MCMC	Markov Chain Monte Carlo

MLM	Modelos Lineares Mistos
MNLM	Modelos Não Lineares Mistos
MOD	Modelo
MV	Máxima Verossimilhança
MVR	Máxima Verossimilhança Restrita
OR	Modelo de Ørskov e McDonald (1979)
RMSE	Raiz do Erro Quadrático Médio
TRV	Teste de Razão de Verossimilhança
UEM	Universidade Estadual de Maringá

SUMÁRIO

Introdução	15
1 Revisão Bibliográfica	17
1.1 Degradabilidade Ruminal <i>in vitro</i>	17
1.1.1 Planejamento temporal e controle de qualidade em ensaios de produção de gases	20
1.1.1.1 Intervalação temporal das leituras e implicações na modelagem cinética	21
1.1.2 Comportamento dos parâmetros cinéticos e ocorrência de valores negativos	22
1.1.3 Degradabilidade Efetiva	24
1.2 Modelos Lineares Mistos	26
1.2.1 Estrutura da matriz de variâncias e covariâncias	30
1.2.1.1 Estruturas de covariância para a matriz de efeitos aleatórios do modelo (D)	31
1.2.1.1.1 Estrutura identidade	31
1.2.1.1.2 Estrutura diagonal heterocedástica	32
1.2.1.1.3 Estrutura de simétrica não estruturada	32
1.2.1.2 Estruturas de covariância para a matriz residual do modelo (R_i)	33
1.2.1.2.1 Estrutura diagonal	33
1.2.1.2.2 Simetria composta (CS)	33
1.2.1.2.3 Estrutura de correlação exponencial (EXP)	34
1.2.1.2.4 Estrutura de correlação Gaussiana (GAUS)	35
1.2.1.3 Estrutura de variância heterogênea (ID)	35
1.3 Modelos Não Lineares Mistos	36
1.3.1 Estimação de β e b	39
1.3.2 Estimação dos componentes de variância	41
1.3.2.1 Máxima Verossimilhança (MV)	41
1.3.2.2 Máxima Verossimilhança Restrita (MVR)	42
1.3.3 Seleção de modelos e qualidade de ajuste	43
1.3.3.1 Teste de Razão de Verossimilhança (TRV)	44
1.3.3.2 Critérios de informação: AIC, AIC _c e BIC	45
1.3.3.3 Intervalos de confiança nos modelos não lineares mistos	46
1.3.3.4 Análise de Resíduos	48

1.3.3.5	Avaliação da capacidade preditiva em modelos não lineares	50
2	Materiais e Métodos	53
2.1	Materiais	53
2.2	Métodos	54
3	Resultados e Discussões	60
3.1	Análise exploratória	60
3.2	Ajuste de modelos não lineares	65
3.3	Ajuste dos modelos não lineares mistos	71
4	Conclusão	86
5	Considerações futuras	88
	Referências	90
	Apêndices	93
APÊNDICE A	Procedimento de truncamento das curvas de produção de gás	94
APÊNDICE B	Estruturas da matriz de variância covariância dos efeitos aleatórios	97
APÊNDICE C	Estimativas dos componentes da matriz de covariância residual R	98
APÊNDICE D	Ajustes adicionais para fins comparativos	99
APÊNDICE E	Resultados completos do teste de Tukey	100

INTRODUÇÃO

A degradabilidade dos alimentos fibrosos no rúmen constitui um processo dinâmico e fundamental para a nutrição de ruminantes, influenciando diretamente a eficiência na utilização de nutrientes, o desempenho zootécnico e os custos de produção animal. Nesse cenário, modelos estatísticos desempenham um papel central ao permitir a descrição da degradação ruminal, a partir da estimativa de parâmetros com significado biológico, como a fração rapidamente degradável, a fração potencialmente degradável, a taxa de degradação e o limite assintótico da digestão.

Dentre as metodologias laboratoriais utilizadas para avaliação da degradabilidade ruminal, destaca-se a técnica de produção de gás *in vitro*, amplamente empregada por sua padronização e capacidade de mensurar a fermentação microbiana por meio da produção acumulada de gases ao longo do tempo. Diferentemente dos métodos *in vivo*, como o uso de animais canulados para a técnica de sacos de nylon, que são considerados procedimentos invasivos e de alta demanda operacional, a técnica *in vitro* se apresenta como uma alternativa não invasiva, com menor custo, maior padronização e controle das condições experimentais (PELL; SCHOFIELD, 1997).

Os dados obtidos por essa metodologia possuem estrutura longitudinal, uma vez que múltiplas medidas são tomadas sequencialmente sobre a mesma unidade experimental, gerando dependência temporal e estrutura hierárquica (LINDSTROM; BATES, 1990). No presente estudo, as observações foram realizadas com intervalos curtos, o que acentua a autocorrelação serial dos resíduos e exige especial atenção à estrutura de covariância.

Apesar da ampla aplicação de modelos não lineares na descrição da dinâmica de produção de gás ruminal, muitos estudos ainda utilizam estruturas simplificadas com efeitos fixos, sem uma avaliação sistemática das diferentes possibilidades de estruturas de variância-covariância para efeitos aleatórios. Os dados obtidos por meio da técnica de produção de gás *in vitro* apresentam duas fontes principais de dependência: a variabilidade entre réplicas experimentais (curvas), associada às diferenças entre unidades experimentais, e a dependência temporal entre observações sucessivas dentro de cada curva, caracterizada por autocorrelação serial.

A seleção da estrutura do modelo é fundamental para inferências fidedignas e interpretações fisiologicamente plausíveis dos parâmetros do modelo. Modelos com bom ajuste estatístico, mas sem adequada fundamentação biológica, podem induzir a conclusões equivocadas. Isso pode levar a implicações na formulação de dietas animais. Além disso, a elevada frequência de medições com curtos intervalos de tempo entre as observações sucessivas intensifica a dependência serial entre os valores registrados, trazendo uma preocupação extra no ajuste dos modelos.

Nesse contexto, os modelos não lineares mistos (MNLMM) constituem uma abordagem estatística robusta e flexível para a análise de dados longitudinais com estrutura de medidas repetidas. Tais modelos integram a modelagem não linear dos efeitos fixos, representando a curva média da população, com a incorporação de efeitos aleatórios, os quais capturam as variações individuais em torno dessa tendência média. Segundo Lindstrom e Bates (1990), os MNLMM combinam as estimativas por mínimos quadrados não lineares para os efeitos fixos com estimativas por máxima verossimilhança (MV) ou máxima verossimilhança restrita (MVR) para os componentes de variância associados aos efeitos aleatórios e ao erro residual, bem como a especificação de diferentes estruturas para a matriz de variância-covariância dos erros intraunidade, possibilitando modelar situações em que há heterocedasticidade ou dependência serial entre as observações dentro da mesma unidade experimental (PINHEIRO; BATES, 2000).

Objetivo: O presente trabalho tem como objetivo ajustar e comparar diferentes modelos não lineares clássicos, entre eles os modelos de Ørskov & McDonald, Gompertz e Logístico, aos dados de produção de gás *in vitro* provenientes de um experimento conduzido com 16 alimentos distintos, incluindo silagens de milho, cana-de-açúcar e gramíneas tropicais, bem como fenos e pré-secados de alfafa, aveia e quicuío. A análise envolve a investigação de diferentes especificações para os efeitos aleatórios associados aos parâmetros dos modelos, bem como a avaliação de distintas estruturas para as matrizes de variância covariância dos efeitos aleatórios e dos erros residuais. Essas especificações permitem representar adequadamente a variabilidade entre unidades experimentais e a dependência intraunidade decorrente das medições repetidas ao longo do tempo. Adicionalmente, são consideradas abordagens baseadas na redefinição dos intervalos temporais entre observações, com o objetivo de atenuar a dependência serial típica de experimentos longitudinais com leituras sucessivas.

A comparação entre os modelos ajustados é conduzida a partir de critérios de informação, como o AIC e o BIC, em conjunto com a análise da log-verossimilhança e a avaliação diagnóstica por meio de gráficos de resíduos, visando compreender o comportamento das diferentes estruturas adotadas e sua influência no ajuste e na capacidade preditiva dos modelos. Nesse contexto, o trabalho busca avaliar de forma sistemática o equilíbrio entre a complexidade das estruturas especificadas e a qualidade do ajuste e estabilidade das estimativas nos modelos considerados.

REVISÃO BIBLIOGRÁFICA

1.1 DEGRADABILIDADE RUMINAL *IN VITRO*

O aumento da demanda por alimentos tem impulsionado o desenvolvimento de técnicas voltadas à melhoria da eficiência produtiva em sistemas pecuários sem comprometer a sustentabilidade dos sistemas de produção. Diante desse cenário, diversas técnicas têm sido desenvolvidas com o objetivo de aumentar a produtividade dos rebanhos e melhorar o aproveitamento dos insumos, assegurando a segurança alimentar (FAO, 2017).

Nesse contexto, a busca por alimentos que promovam maior desempenho animal, especialmente em ruminantes, tem se intensificado. A eficiência alimentar torna-se, portanto, um fator crucial, não apenas sob o ponto de vista econômico, mas também ambiental. Entre as diferentes abordagens voltadas à compreensão e ao aprimoramento desse processo, destaca-se a avaliação da degradabilidade ruminal *in vitro*, uma metodologia laboratorial que permite estimar o potencial fermentativo dos alimentos em ambiente simulado.

Segundo Filho et al. (2006), a técnica *in vitro* é fundamental para a estimativa da degradabilidade da matéria seca e da proteína bruta dos alimentos e das frações fibrosas, fornecendo dados que subsidiam a formulação de dietas mais eficientes para ruminantes. A aplicação de modelos matemáticos, como o proposto por Ørskov e McDonald (1979), permite descrever a cinética da degradação de nutrientes, considerando frações solúveis, potencialmente degradáveis e indigestíveis, bem como as respectivas taxas de degradação e passagem ruminal.

A técnica baseia-se na incubação de amostras de alimentos em frascos contendo fluido ruminal fresco, tampão e meio de cultura, sob condições rigorosamente controladas de temperatura (39 °C), pH e anaerobiose. De acordo com Filho et al. (2006), essas condições experimentais permitem estimar a fração degradada ao longo do tempo, refletindo a dinâmica fermentativa dos substratos no rúmen.

Entre as variantes da técnica destaca-se a produção de gases *in vitro*, a qual quantifica o volume de gases gerados pelo metabolismo microbiano durante a fermentação dos substratos. Tal volume apresenta elevada correlação com a digestibilidade e a degradabilidade da matéria

orgânica, configurando-se como indicador eficiente do potencial fermentativo dos alimentos. A medição periódica desses volumes permite a construção de curvas cinéticas de degradação, que podem ser ajustadas por modelos não lineares e interpretadas por meio de parâmetros biológicos como o volume total de gás, a taxa fracional de fermentação e o tempo de colonização microbiana (FRANCE; DIJKSTRA, 2005).

Como destaca Mertens (2016), os parâmetros derivados dessas curvas descrevem não apenas o padrão fermentativo, mas também aspectos intrínsecos dos alimentos relacionados à sua disponibilidade efetiva para os microrganismos ruminais. Nesse sentido, os modelos matemáticos aplicados à produção de gases assumem papel central na predição da eficiência alimentar, contribuindo para o planejamento de experimentos e para a formulação de dietas com melhor aproveitamento biológico e menor impacto ambiental.

A crescente sofisticação das técnicas de incubação *in vitro* e a complexidade dos dados gerados, especialmente medidas repetidas ao longo do tempo, demandam abordagens estatísticas capazes de modelar tanto a tendência média da degradação quanto a variabilidade entre unidades experimentais. Nesse cenário, os modelos mistos não lineares destacam-se como ferramentas robustas e versáteis para o ajuste de curvas de degradação.

Esses modelos permitem a incorporação simultânea de efeitos fixos, responsáveis por descrever os fatores sistemáticos (como tipo de alimento ou tempo de incubação), e de efeitos aleatórios, que capturam a variabilidade intra e interindividual entre frascos, animais doadores ou repetições. Essa estrutura hierárquica é especialmente útil em ensaios de produção de gás, nos quais há repetidas medições ao longo do tempo em diferentes unidades experimentais (PINHEIRO; BATES, 2000).

A combinação entre a não linearidade das equações biológicas e a modelagem da estrutura de correlação dos dados longitudinais confere aos modelos não lineares mistos maior realismo na estimativa dos parâmetros cinéticos, além de permitir inferências mais precisas e generalizáveis. Conforme apontado por Pinheiro e Bates (2000), tais modelos são particularmente apropriados em contextos nos quais a resposta depende de funções exponenciais ou sigmóides, como ocorre com a degradação de nutrientes no rúmen.

Nesse contexto, o uso de modelos não lineares mistos na análise de dados de degradabilidade ruminal constitui uma abordagem metodológica particularmente adequada, pois permite conciliar a natureza intrinsecamente não linear do processo fermentativo com a estrutura longitudinal das observações ao longo do tempo. Além de assegurar maior rigor estatístico na estimação dos parâmetros, essa abordagem possibilita uma representação mais fiel da variabilidade biológica entre unidades experimentais, resultando em estimativas mais robustas e biologicamente interpretáveis.

A escolha dos modelos de Ørskov e McDonald, Gompertz e Logístico fundamenta-se na ampla utilização dessas funções na literatura para descrever a cinética de degradação ruminal e

de produção de gás, bem como na capacidade dessas equações em representar diferentes padrões de resposta ao longo do tempo. Conforme discutido por Teixeira et al. (2016), esses modelos diferem em sua estrutura funcional, mas compartilham a possibilidade de decompor o processo fermentativo em componentes associados à fração rapidamente disponível, ao potencial total de degradação e à velocidade com que esse potencial é alcançado.

No modelo de Ørskov e McDonald, originalmente proposto para dados de degradabilidade ruminal, a função exponencial descreve um comportamento assintótico no qual a resposta aumenta rapidamente nos instantes iniciais e tende a um valor máximo ao longo do tempo. Nessa formulação, os parâmetros estão diretamente associados à fração solúvel do alimento, à fração insolúvel potencialmente degradável e à taxa fracional de degradação dessa fração, sendo essa interpretação amplamente aceita e empregada como referência em estudos de cinética ruminal.

Por sua vez, os modelos Gompertz e Logístico apresentam natureza sigmoide, sendo mais flexíveis para descrever processos nos quais a resposta fermentativa não se inicia imediatamente de forma exponencial, mas apresenta uma fase inicial de menor atividade microbiana, seguida por um período de rápida intensificação da produção de gás e, posteriormente, por uma estabilização assintótica. De acordo com Teixeira et al. (2016), essas funções são particularmente adequadas para descrever variações na dinâmica da degradação quando a resposta não é bem representada por um crescimento exponencial simples.

Para todos os modelos considerados neste trabalho utilizou-se uma forma reparametrizada em que os parâmetros apresentassem interpretações comparáveis entre as diferentes funções, conforme estratégia utilizada por Teixeira et al. (2016). Essa reparametrização permite que os modelos sejam comparados tanto do ponto de vista estatístico quanto biológico, sem implicar equivalência estrita entre suas estruturas funcionais.

As expressões matemáticas consideradas para a i -ésima réplica ($i = 1, \dots, n$) da j -ésima amostra ($j = 1, \dots, m$), observada no tempo t_{ij} , são apresentadas a seguir.

- Modelo de Ørskov e McDonald:

$$y_{ij} = \beta_1 + \beta_2 \left[1 - \exp(-\beta_3 t_{ij}) \right] + \varepsilon_{ij}. \quad (1.1)$$

- Modelo de Gompertz:

$$y_{ij} = \beta_1 + \beta_2 \exp \left[-\exp(1 - \beta_3 t_{ij}) \right] + \varepsilon_{ij}. \quad (1.2)$$

- Modelo Logístico:

$$y_{ij} = \beta_1 + \beta_2 / \left[1 + \exp(2 - 4\beta_3 t_{ij}) \right] + \varepsilon_{ij}. \quad (1.3)$$

Em todas as formulações, y_{ij} representa o volume acumulado de gás observado no tempo t_{ij} e ε_{ij} denota o erro aleatório associado à observação, assumido com distribuição normal de média zero.

De forma geral, os parâmetros β_1 , β_2 e β_3 estão associados, respectivamente, à fração solúvel (%), à fração insolúvel potencialmente degradável (%) e à taxa fracional de degradação. A utilização de uma nomenclatura paramétrica comum aos diferentes modelos, conforme adotado por Teixeira et al. (2016), permite uma comparação conjunta do desempenho estatístico e da coerência biológica das funções ajustadas.

Embora os parâmetros tenham interpretações análogas, suas implicações biológicas não são estritamente equivalentes, devido às diferenças estruturais entre as funções. Assim, a adoção de uma nomenclatura paramétrica comum não implica equivalência biológica estrita entre os modelos, mas constitui uma estratégia metodológica que facilita a comparação conjunta da qualidade de ajuste e do significado biológico dos parâmetros estimados.

A análise comparativa entre os três modelos, tanto em sua forma com efeitos fixos quanto na estrutura mista, visa identificar não apenas o modelo com melhor desempenho, mas também aquele que apresente coerência biológica na descrição da cinética fermentativa. A escolha de modelos para dados de produção de gás deve considerar simultaneamente a qualidade do ajuste e a interpretação fisiológica dos parâmetros estimados, uma vez que modelos com bom desempenho estatístico, mas sem fundamento biológico, podem gerar inferências equivocadas (PELL; SCHOFIELD, 1997). Além disso, Schofield e Pell (1994) ressaltam que modelos sigmóides como Gompertz e Logístico tendem a ajustar melhor a dados oriundos de sistemas automatizados de incubação, devido à capacidade de capturar as fases distintas da fermentação microbiana.

Complementarmente, a utilização de modelos mistos com diferentes estruturas de variância-covariância para os efeitos aleatórios, como identidade, diagonal heterocedástica e de simetria não estruturada, permite representar distintos padrões de variabilidade entre as unidades experimentais (PINHEIRO; BATES, 2000). A escolha adequada dessa estrutura é essencial para garantir estimativas confiáveis, pois erros de especificação podem afetar tanto a significância estatística dos efeitos fixos quanto a interpretação dos parâmetros do modelo (VERBEKE; MOLENBERGHS, 2009). Nesse contexto, a comparação entre estruturas visa selecionar aquela que melhor descreve a variabilidade interindividual presente nos experimentos de produção de gás *in vitro*.

1.1.1 PLANEJAMENTO TEMPORAL E CONTROLE DE QUALIDADE EM ENSAIOS DE PRODUÇÃO DE GASES

Ensaio *in vitro* de produção de gases requerem planejamento cuidadoso tanto do esquema temporal de coleta quanto do controle de qualidade das curvas obtidas, uma vez que a dinâmica fermentativa é fortemente dependente das condições experimentais e da consistência das medições ao longo do tempo. Protocolos clássicos destacam a necessidade de manutenção rigorosa da temperatura e do pH, uso de frascos controle (brancos) para correção da produção

basal de gás e vedação adequada dos frascos, de modo a garantir que a produção observada seja efetivamente atribuída à fermentação do substrato incubado (THEODOROU et al., 1994).

Adicionalmente, realiza-se uma etapa de controle de qualidade baseada na inspeção visual das curvas de produção de gás, com o objetivo de identificar comportamentos inconsistentes, como quedas na curva acumulada, saltos abruptos ou padrões incompatíveis com a dinâmica fermentativa esperada. Nesses casos, procede-se ao truncamento das curvas ou à exclusão de observações a partir de tempos previamente definidos para cada unidade experimental, de modo a evitar a influência de artefatos experimentais no ajuste dos modelos. Essa estratégia permite maior estabilidade na estimação dos parâmetros e maior coerência biológica das curvas ajustadas.

Além disso, a escolha dos tempos de incubação exerce papel central na identificação adequada das diferentes fases da fermentação, incluindo o período inicial, a fase de maior atividade microbiana e a aproximação assintótica do processo. A literatura enfatiza que esquemas temporais mal distribuídos ou excessivamente curtos podem mascarar a verdadeira dinâmica do processo, dificultando tanto a interpretação biológica quanto o ajuste estatístico dos modelos (FRANCE; THORNLEY, 1984).

1.1.1.1 INTERVALAÇÃO TEMPORAL DAS LEITURAS E IMPLICAÇÕES NA MODELAGEM CINÉTICA

A escolha da intervalação temporal das leituras em ensaios de produção de gases *in vitro* exerce influência direta sobre a qualidade da descrição da cinética fermentativa e sobre a precisão na estimação dos parâmetros dos modelos matemáticos empregados. A literatura evidencia que a produção cumulativa de gases apresenta comportamento não linear, caracterizado por rápida variação nas fases iniciais da incubação e posterior desaceleração até a aproximação assintótica em tempos mais longos (ØRSKOV; MCDONALD, 1979).

Teixeira et al. (2016) destacam que esquemas de amostragem temporal devem priorizar maior densidade de observações nas primeiras horas de incubação, período crítico para a adequada identificação das taxas de degradação e do possível tempo de latência, especialmente em modelos exponenciais e sigmoides. Protocolos que empregam sistemas automatizados de mensuração de gases têm adotado leituras em intervalos inferiores a uma hora nas fases iniciais, justamente com o objetivo de capturar a dinâmica exponencial da fermentação e reduzir a imprecisão na estimação da taxa (THEODOROU et al., 1994). A ausência de leituras nesse intervalo pode comprometer a identificabilidade desses parâmetros, aumentando a variância das estimativas e a correlação entre componentes do vetor paramétrico.

Por outro lado, estudos recentes enfatizam que a inclusão de tempos prolongados de incubação é fundamental para a correta caracterização da fração indigestível dos alimentos, em particular da fibra detergente neutro indigestível (uNDF). Trabalhos voltados à modelagem de múltiplos compartimentos e à formulação de rações ressaltam que leituras tardias são essenciais

para a identificação do comportamento assintótico da fermentação e para a separação adequada entre frações potencialmente degradáveis e não degradáveis (LI et al., 2023; GU et al., 2024).

Assim, embora o esquema de intervalação temporal seja previamente definido no delineamento experimental, sua adequação é avaliada à luz de aspectos biológicos e estatísticos inerentes à modelagem adotada. De modo geral, esquemas temporais empregados em experimentos de produção de gás *in vitro* apresentam resolução suficiente nas fases iniciais, sem excessiva densidade de pontos, buscando equilibrar a captura da dinâmica fermentativa com a necessidade de evitar elevada dependência temporal entre observações sucessivas. Tal abordagem está em consonância com protocolos experimentais consolidados na literatura, nos quais a definição dos tempos de leitura busca contemplar simultaneamente a fase inicial de rápida fermentação e a estabilização assintótica do processo como o apresentado por Cabral et al. (2019).

Em modelos longitudinais, observações muito próximas no tempo tendem a apresentar alta correlação residual, o que pode resultar em estruturas de covariância próximas da singularidade e dificuldades na estimação por máxima verossimilhança. Tal situação torna-se particularmente relevante quando se consideram modelos não lineares mistos com múltiplos níveis de variabilidade, como no presente estudo, que envolve 16 tratamentos, três parâmetros por função e efeitos aleatórios associados a cada um desses parâmetros. O aumento excessivo no número de observações por unidade experimental poderia intensificar a dependência temporal e ampliar a complexidade da estrutura de covariância estimada, comprometendo a estabilidade numérica e a convergência do algoritmo (ARCHONTOULIS; MIGUEZ, 2015).

Dessa forma, o esquema temporal adotado buscou manter quantidade adequada de observações para garantir identificabilidade e precisão dos parâmetros, sem introduzir redundância excessiva de informação temporal. Essa estratégia contribui para melhor condicionamento da matriz de informação e para maior robustez na estimação simultânea dos efeitos fixos e aleatórios, especialmente em contextos de modelagem com múltiplos tratamentos e estrutura hierárquica de variabilidade.

1.1.2 COMPORTAMENTO DOS PARÂMETROS CINÉTICOS E OCORRÊNCIA DE VALORES NEGATIVOS

Durante o ajuste de modelos não lineares a dados de degradabilidade ruminal, a ocorrência de estimativas de parâmetros com valores negativos ou biologicamente implausíveis é relativamente frequente, sobretudo em experimentos nos quais há elevada variabilidade entre unidades experimentais, limitação de observações nos tempos iniciais de incubação ou dificuldades de convergência numérica dos métodos iterativos de estimação. Conforme discutido por Sartorio (2013) em sua tese e apresentado em Medeiros et al. (2020), tais ocorrências não devem ser interpretadas, a priori, como falhas do modelo matemático, mas como indícios de incompatibilidade entre a estrutura funcional do modelo, a qualidade dos dados observados e os

fenômenos biológicos subjacentes ao processo fermentativo.

No contexto do modelo clássico de Ørskov e McDonald (1979), parâmetros negativos associados à fração rapidamente solúvel ou ao potencial de degradação podem emergir quando o alimento apresenta um período pré-fermentativo, durante o qual não se observa degradação efetiva do substrato. Esse intervalo, compreendido entre o início da incubação e o início da atividade microbiana efetiva, é denominado tempo de latência ou tempo de colonização (*lag time*), e está relacionado a processos como hidratação das partículas, remoção de substâncias inibidoras e adesão dos microrganismos ruminais às superfícies do alimento (SAVIAN, 2008).

Quando o *lag time* não é explicitamente incorporado ao modelo, parte desse efeito pode ser absorvida pelos parâmetros estruturais da curva, resultando em estimativas negativas, particularmente para o parâmetro associado à fração imediatamente solúvel. Segundo Thiago (1994), valores negativos desse parâmetro indicam a necessidade de um tempo prévio para o início da degradação, sugerindo que a fração solúvel não é prontamente disponível no instante inicial da incubação. Essa interpretação é especialmente relevante na avaliação de forragens tropicais, nas quais o processo de lignificação tende a retardar a colonização microbiana.

Entretanto, Mertens (1993) alerta que, em determinadas situações, o próprio tempo de colonização estimado pode assumir valores negativos, o que compromete sua interpretação biológica. Um *lag time* negativo implicaria que a digestão se inicia antes do tempo zero, resultado considerado biologicamente inconsistente. Segundo Mertens e Ely (1982), tal comportamento decorre, em geral, de uma solubilização praticamente instantânea do substrato, aliada à ausência de observações suficientemente precoces para permitir a identificação adequada dessa fase inicial. Nessas circunstâncias, erros na estimação das taxas fracionais de digestão e das frações potencialmente degradáveis e indigestíveis tornam-se mais prováveis.

Diante dessas limitações, a literatura recomenda cautela na interpretação de parâmetros negativos e desaconselha a imposição indiscriminada de restrições apenas para forçar plausibilidade biológica. Conforme enfatizado por Mertens (1993) e retomado por Medeiros et al. (2020), a presença de valores negativos deve motivar uma análise criteriosa do traçado das curvas individuais, da adequação do esquema temporal de incubação e da escolha do modelo cinético, incluindo a possibilidade de reparametrizações ou da adoção de modelos que considerem explicitamente o tempo de colonização.

Em estudos baseados na técnica de produção de gases *in vitro*, situação análoga pode ser observada quando a fase inicial de fermentação não é adequadamente compreendida. Nesses casos, parâmetros negativos associados à produção acumulada de gás ou às taxas de fermentação refletem, predominantemente, limitações do desenho experimental e da especificação do modelo, e não comportamentos biologicamente impossíveis. Assim, a interpretação desses resultados deve ser conduzida à luz do processo fermentativo e das condições experimentais, priorizando a coerência biológica e a estabilidade do ajuste em detrimento de restrições artificiais aos parâmetros.

1.1.3 DEGRADABILIDADE EFETIVA

No âmbito da avaliação nutricional, a degradabilidade efetiva (DE) constitui uma medida sintética que integra informações cinéticas do processo fermentativo, permitindo quantificar a fração do substrato potencialmente degradável que seria efetivamente utilizada no rúmen ao longo do tempo. Conceitualmente, a DE decorre da decomposição em frações e taxas de degradação e é tradicionalmente definida no contexto de ensaios *in situ*, nos quais se considera explicitamente uma taxa de passagem ruminal (ØRSKOV; MCDONALD, 1979).

Em ensaios *in vitro*, a produção acumulada de gás é interpretada como um marcador indireto da atividade fermentativa microbiana e, por consequência, do potencial de degradação do substrato incubado. Assim, a utilização de medidas análogas à DE em dados *in vitro* deve ser compreendida como uma métrica comparativa do potencial fermentativo efetivo entre tratamentos, e não como uma estimativa direta de digestibilidade *in vivo*. A degradabilidade efetiva também pode ser obtida a partir de ensaios *in vitro* de produção de gases, por meio do ajuste de modelos cinéticos à fermentação e da incorporação de uma taxa de passagem aos parâmetros estimados (KAMALAK et al., 2005).

A DE é usualmente calculada a partir dos parâmetros estimados no modelo cinético, considerando a fração potencialmente degradável (β_2) e a taxa de degradação (β_3), ponderadas por uma taxa de passagem ruminal fixa. Para o modelo de Ørskov e McDonald (1979), cuja formulação foi apresentada na Equação (1.1), a DE pode ser expressa como uma função dos parâmetros do modelo, dada por:

$$DE = \beta_1 + \frac{\beta_2 \cdot \beta_3}{\beta_3 + k}, \quad (1.4)$$

em que k representa taxa de passagem ruminal, considerada fixa. Essa expressão decorre diretamente da parametrização do modelo de Ørskov e McDonald, no qual os parâmetros β_1 , β_2 e β_3 são estimados a partir do ajuste da curva de produção de gás descrita na Equação (1.1). Teixeira et al. (2016) avaliaram diferentes funções não lineares, incluindo os modelos Gompertz e Logístico, para descrever a cinética de produção de gases *in vitro*. Embora a expressão clássica da degradabilidade efetiva esteja formalmente associada ao modelo exponencial de Ørskov e McDonald, o autor destaca que medidas integradas da cinética fermentativa podem ser obtidas a partir de outras parametrizações, desde que preservada a interpretação biológica dos parâmetros.

Como a degradabilidade efetiva (DE) é definida como função dos parâmetros estimados no modelo cinético, sua inferência deve considerar a incerteza associada à estimação desses parâmetros. Para esse fim, utilizou-se o método delta, amplamente empregado para aproximar a variância de funções diferenciáveis de estimadores obtidos por máxima verossimilhança (CASELLA; BERGER, 2002).

Para cada tratamento, a DE foi expressa como função do vetor de parâmetros estimados, considerando a parametrização específica do modelo ajustado. A variância aproximada foi calculada com base na matriz de variância covariância dos efeitos fixos fornecida no processo

de estimação do modelo misto (PINHEIRO; BATES, 2000). Formalmente,

$$\text{Var}(\widehat{\text{DE}}) \approx \nabla g(\widehat{\beta})^\top \widehat{V}_\beta \nabla g(\widehat{\beta}), \quad (1.5)$$

sendo $\widehat{\beta} = (\widehat{\beta}_1, \widehat{\beta}_2, \widehat{\beta}_3)^\top$ o vetor de estimativas dos parâmetros do modelo, $\widehat{V}_\beta$ a matriz de variância covariância dos estimadores dos efeitos fixos e $\nabla g(\widehat{\beta})$ corresponde ao gradiente da função que define a DE, avaliado nas estimativas dos parâmetros. Os intervalos de confiança de 95% foram obtidos assumindo aproximação normal, isto é,

$$\text{IC}_{95\%} = \widehat{\text{DE}} \pm 1,96 \times \sqrt{\text{Var}(\widehat{\text{DE}})}, \quad (1.6)$$

Em ensaios *in vitro* de produção de gases, a taxa de passagem ruminal (k) não é observada experimentalmente, sendo, portanto, tratada como um parâmetro externo ao modelo. A adoção de valores fixos de k (por exemplo, 0,02; 0,05; 0,08 h⁻¹) baseia-se em recomendações clássicas da literatura, com o objetivo de padronizar comparações entre alimentos e tratamentos. O valor de $k = 0,05 \text{ h}^{-1}$ tem sido amplamente utilizado para representar condições médias de consumo e taxa de passagem, sendo considerado adequado em estudos comparativos de cinética de degradação e produção de gases (TOMICH et al., 2006).

Resultados empíricos reforçam essa abordagem ao indicar que, embora os parâmetros cinéticos estimados por metodologias *in situ* e *in vitro* possam diferir substancialmente, especialmente no que se refere à taxa de degradação, as estimativas de DE tendem a apresentar valores próximos entre técnicas, por se tratar de uma medida integrada da cinética fermentativa (TOMICH et al., 2006). Nesse sentido, López et al. (1999) destacam que a utilização de parâmetros integrados constitui uma abordagem adequada para a comparação do potencial fermentativo relativo de substratos avaliados por curvas de produção de gases *in vitro*, desde que tais medidas sejam interpretadas com cautela.

Sendo assim, a avaliação da degradabilidade ruminal *in vitro* por meio da produção de gases envolve decisões que vão além da escolha do protocolo experimental, abrangendo também a especificação dos modelos matemáticos e das estruturas estatísticas utilizadas na análise dos dados. Os aspectos relacionados à cinética fermentativa, à intervalação temporal das leituras, ao comportamento dos parâmetros e à ocorrência de estimativas biologicamente implausíveis fornecem a base conceitual para as escolhas metodológicas adotadas neste estudo. Nas seções seguintes, são apresentados os materiais, o delineamento experimental e os procedimentos estatísticos empregados, bem como as estratégias de ajuste, comparação e avaliação dos modelos não lineares mistos considerados.

1.2 MODELOS LINEARES MISTOS

Em experimentos que envolvem observações repetidas sobre a mesma unidade experimental ao longo do tempo, como nos ensaios de degradabilidade ruminal, os dados apresentam natureza longitudinal. Nessas situações, as medições realizadas na mesma unidade experimental tendem a ser correlacionadas, podendo ainda ocorrer heterogeneidade de variâncias entre os diferentes instantes de avaliação. A não consideração explícita desses aspectos pode comprometer tanto as estimativas dos parâmetros quanto os erros padrão associados a eles (SARTORIO, 2013).

A análise desses dados tem sido tradicionalmente realizada por meio da abordagem univariada de perfis (ANOVA) ou pela análise multivariada de perfis (MANOVA), ambas fundamentadas em hipóteses restritivas acerca da estrutura da matriz de variâncias–covariâncias entre os tempos de observação.

Na abordagem univariada de perfis, frequentemente associada a um delineamento em parcelas subdivididas (*split-plot design*), impõe-se uma estrutura específica à matriz de variâncias–covariâncias entre os instantes de avaliação. Nessa formulação, admite-se que a variância das respostas seja constante ao longo do tempo e que a covariância entre quaisquer dois instantes seja igual para todos os pares de medidas realizadas na mesma unidade experimental.

Essa estrutura pode ser representada por:

$$\Sigma = \begin{bmatrix} \sigma_\varepsilon^2 + \sigma_\gamma^2 & \sigma_\gamma^2 & \cdots & \sigma_\gamma^2 \\ \sigma_\gamma^2 & \sigma_\varepsilon^2 + \sigma_\gamma^2 & \cdots & \sigma_\gamma^2 \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_\gamma^2 & \sigma_\gamma^2 & \cdots & \sigma_\varepsilon^2 + \sigma_\gamma^2 \end{bmatrix}, \quad (1.7)$$

em que σ_ε^2 representa a variância residual e σ_γ^2 corresponde à covariância comum entre quaisquer dois tempos avaliados na mesma unidade experimental.

Tal restrição pode ser inadequada para dados de degradabilidade ruminal, nos quais é esperado que a correlação entre observações diminua à medida que a distância temporal aumenta, além da possibilidade de variâncias distintas entre fases iniciais e finais da incubação.

Como alternativa, a MANOVA admite uma matriz de variâncias–covariâncias não estruturada, permitindo que cada tempo possua variância própria e que cada par de tempos tenha sua própria covariância. Para t tempos de medição, essa estrutura envolve $\frac{t(t+1)}{2}$ parâmetros livres. Embora mais flexível, tal abordagem pode tornar-se instável à medida que t cresce, especialmente em amostras pequenas, além de exigir dados completos e pressupor normalidade multivariada (WEST; WELCH; GALECKI, 2014).

Essa estrutura é representada pela seguinte matriz Σ , com t tempos de medição como exemplo:

$$\Sigma = \begin{bmatrix} \sigma_1^2 & \sigma_{12} & \cdots & \sigma_{1t} \\ \sigma_{21} & \sigma_2^2 & \cdots & \sigma_{2t} \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_{t1} & \sigma_{t2} & \cdots & \sigma_t^2 \end{bmatrix}. \quad (1.8)$$

Conforme discutido por Sartorio (2013) e apresentado em Medeiros et al. (2020), nas análises univariadas de perfis os testes envolvendo comparações intraindivíduos (associadas ao fator tempo) são exatos apenas quando a matriz de variâncias–covariâncias satisfaz a condição de esfericidade. Caso essa condição seja rejeitada, a análise univariada do tipo *split-plot* não deve ser empregada.

Como alternativa, pode-se recorrer à metodologia multivariada (MANOVA), que admite uma estrutura de variância covariância menos restritiva. Entretanto, embora essa abordagem não imponha a mesma condição estrutural da análise univariada, ela ainda não se mostra totalmente adequada para representar a estrutura de covariâncias de dados longitudinais.

Adicionalmente, a MANOVA não incorpora com facilidade covariáveis que variam ao longo do tempo, desconsidera dados coletados em tempos não regulares e, assim como o método univariado, apresenta dificuldades na presença de estruturas desbalanceadas (SARTORIO, 2013).

Como alternativa, os modelos lineares mistos (MLM) permitem representar simultaneamente a estrutura hierárquica dos dados e a dependência entre observações repetidas. Nessa abordagem, incluem-se efeitos fixos, associados a fatores como tempo e tratamento, e efeitos aleatórios, relacionados às unidades experimentais ou repetições, possibilitando modelar a variabilidade entre e dentro das unidades amostrais.

Do ponto de vista estrutural, os modelos mistos permitem especificar explicitamente a matriz de variância covariância das observações, seja por meio da modelagem dos efeitos aleatórios, seja pela escolha de estruturas residuais adequadas, cuja seleção pode ser orientada por critérios de informação como AIC, BIC e log-verossimilhança.

O desenvolvimento formal dessa abordagem remonta aos trabalhos pioneiros de Henderson (1950), que estabeleceram as equações do modelo misto para estimação conjunta de efeitos fixos e aleatórios. Posteriormente, a formulação foi ampliada e difundida para dados longitudinais por autores como Laird e Ware (1982) e sistematizada em obras de referência como Verbeke e Molenberghs (2009) e Pinheiro e Bates (2000), consolidando-se como ferramenta fundamental para a análise de dados com medidas repetidas.

Modelos lineares de efeitos mistos, também conhecidos como modelos mistos, são uma generalização dos modelos lineares tradicionais, incorporando, além dos efeitos fixos, componentes aleatórios que representam fontes de variação associadas a unidades amostrais específicas ou níveis hierárquicos do experimento. De acordo com Pinheiro e Bates (2000),

essa estrutura possibilita a modelagem de dados correlacionados, como aqueles provenientes de medidas repetidas, e permite estimativas mais eficientes dos parâmetros quando comparada às abordagens clássicas.

A formulação de modelos mistos pode ser apresentada de forma hierárquica, descrevendo explicitamente os componentes intra e interindividuais. Essa abordagem é particularmente útil para experimentos com medidas repetidas, como no caso da cinética de produção de gás em forragens, em que múltiplas observações são realizadas em cada unidade experimental ao longo do tempo.

Conforme apresentado em Laird e Ware (1982), para cada unidade experimental i , com $i = 1, \dots, n$, seja $\mathbf{y}_i = (y_{i1}, y_{i2}, \dots, y_{in_i})^\top$ o vetor de respostas observadas em n_i tempos distintos. Considera-se inicialmente o modelo de regressão individual:

$$\mathbf{y}_i = \mathbf{Z}_i \boldsymbol{\beta}_i + \boldsymbol{\varepsilon}_i, \quad (1.9)$$

em que $\mathbf{Z}_i$ é uma matriz de delineamento ($n_i \times q$) de covariáveis, $\boldsymbol{\beta}_i$ é o vetor de coeficientes de regressão específicos para o indivíduo i , e $\boldsymbol{\varepsilon}_i$ é o vetor de erros aleatórios intraindividuais, com distribuição $\boldsymbol{\varepsilon}_i \sim N_{n_i}(\mathbf{0}, \mathbf{R}_i)$, sendo $\mathbf{R}_i$ a matriz de covariância residual, que reflete a estrutura de correlação entre as medidas ao longo do tempo.

Para descrever a variabilidade entre indivíduos, assume-se que os vetores de coeficientes $\boldsymbol{\beta}_i$ são decompostos em uma parte fixa comum a todos os indivíduos e uma parte aleatória:

$$\boldsymbol{\beta}_i = \mathbf{A}_i \boldsymbol{\beta} + \mathbf{b}_i, \quad (1.10)$$

em que $\mathbf{A}_i$ é uma matriz de delineamento ($q \times p$) para os efeitos fixos, $\boldsymbol{\beta}$ é o vetor de parâmetros fixos e $\mathbf{b}_i$ é um vetor de efeitos aleatórios associados ao indivíduo i , com $\mathbf{b}_i \sim N_q(0, \mathbf{D})$, sendo $\mathbf{D}$ a matriz de variância e covariância dos efeitos aleatórios. Assume-se ainda independência entre $\boldsymbol{\varepsilon}_i$ e $\mathbf{b}_i$.

Substituindo a equação (1.10) na equação (1.9), obtém-se a formulação do modelo linear misto para o i -ésimo indivíduo:

$$\mathbf{y}_i = \mathbf{X}_i \boldsymbol{\beta} + \mathbf{Z}_i \mathbf{b}_i + \boldsymbol{\varepsilon}_i, \quad (1.11)$$

sendo $\mathbf{X}_i = \mathbf{Z}_i \mathbf{A}_i$ a matriz de delineamento associada aos efeitos fixos, com:

$$\mathbf{y}_i = \begin{bmatrix} y_{i1} \\ y_{i2} \\ \vdots \\ y_{in_i} \end{bmatrix}, \quad \mathbf{X}_i = \begin{bmatrix} x_{i1}^{(1)} & x_{i1}^{(2)} & \cdots & x_{i1}^{(p)} \\ x_{i2}^{(1)} & x_{i2}^{(2)} & \cdots & x_{i2}^{(p)} \\ \vdots & \vdots & \ddots & \vdots \\ x_{in_i}^{(1)} & x_{in_i}^{(2)} & \cdots & x_{in_i}^{(p)} \end{bmatrix}, \quad \boldsymbol{\beta} = \begin{bmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_q \end{bmatrix},$$

$$\mathbf{Z}_i = \begin{bmatrix} z_{i1}^{(1)} & z_{i1}^{(2)} & \cdots & z_{i1}^{(q)} \\ z_{i2}^{(1)} & z_{i2}^{(2)} & \cdots & z_{i2}^{(q)} \\ \vdots & \vdots & \ddots & \vdots \\ z_{in_i}^{(1)} & z_{in_i}^{(2)} & \cdots & z_{in_i}^{(q)} \end{bmatrix}, \quad \mathbf{b}_i = \begin{bmatrix} b_{i1} \\ b_{i2} \\ \vdots \\ b_{iq} \end{bmatrix}, \quad \boldsymbol{\varepsilon}_i = \begin{bmatrix} \varepsilon_{i1} \\ \varepsilon_{i2} \\ \vdots \\ \varepsilon_{in_i} \end{bmatrix}.$$

Supondo $\mathbf{b}_i \sim N(0, \mathbf{D})$ e $\boldsymbol{\varepsilon}_i \sim N(0, \mathbf{R}_i)$, e assumindo independência entre os termos de erro e os efeitos aleatórios, a distribuição marginal do vetor de respostas $\mathbf{y}_i = \mathbf{X}_i \boldsymbol{\beta} + \boldsymbol{\varepsilon}_i^*$ é dada por:

$$\mathbf{y}_i \sim N(\mathbf{X}_i \boldsymbol{\beta}, \boldsymbol{\Sigma}_i), \quad (1.12)$$

em que a matriz de variância-covariância marginal é expressa por $\boldsymbol{\Sigma}_i = \mathbf{Z}_i \mathbf{D} \mathbf{Z}_i^\top + \mathbf{R}_i$ e $\boldsymbol{\varepsilon}_i^*$ representa o erro residual marginal com $\boldsymbol{\varepsilon}_i^* \sim N(0, \boldsymbol{\Sigma}_i^*)$. Conforme apresentado por West, Welch e Galecki (2014), a média marginal (ou valor esperado) e a matriz de variância covariância marginal do vetor de respostas $\mathbf{y}_i$ são dadas, respectivamente, por

$$\mathbb{E}(\mathbf{y}_i) = \mathbf{X}_i \boldsymbol{\beta}, \quad (1.13)$$

e

$$\text{Var}(\mathbf{y}_i) = \boldsymbol{\Sigma}_i = \mathbf{Z}_i \mathbf{D} \mathbf{Z}_i^\top + \mathbf{R}_i, \quad (1.14)$$

Essa formulação permite descrever simultaneamente a variabilidade entre a unidade experimental observada (por meio de $\mathbf{D}$) e a dependência das medidas repetidas dentro de cada unidade experimental (via $\mathbf{R}_i$).

A formulação hierárquica dos modelos mistos permite não apenas estimar parâmetros fixos e aleatórios com robustez, mas também testar diferentes estruturas de variância-covariância para os resíduos, como simetria composta, autoregressiva ou estrutura não estruturada (PINHEIRO; BATES, 2000).

No entanto, quando a relação entre as variáveis preditoras (como o tempo de incubação) e a variável resposta (volume de gás, por exemplo) não é linear, como é frequentemente observado em processos biológicos que envolvem curvas com padrões sigmóides ou exponenciais, a utilização de modelos mistos lineares torna-se limitada. Nesses casos, a não linearidade estrutural do modelo precisa ser incorporada diretamente à especificação matemática do modelo.

Conforme destacado por Pinheiro e Bates (2000), os modelos não lineares mistos (MNLM) oferecem a flexibilidade necessária para ajustar funções com significado biológico, acomodando variações entre e dentro das unidades amostrais e preservando a interpretação fisiológica dos parâmetros estimados. Os autores ressaltam que essa abordagem é particularmente vantajosa em estudos nos quais a curva média da população segue um comportamento não linear, e deseja-se também captar as diferenças individuais em relação a essa curva.

Nesse mesmo sentido, West, Welch e Galecki (2014) enfatizam que modelos não lineares mistos são a escolha mais apropriada quando os dados apresentam dependência temporal e a

função que descreve a resposta depende de parâmetros que influenciam a forma da curva, como assíntotas, taxas de crescimento ou pontos de inflexão. Tais características são típicas de experimentos com fermentação ruminal, onde os padrões de produção de gás ao longo do tempo apresentam comportamentos sigmóides que não podem ser representados por funções lineares.

1.2.1 ESTRUTURA DA MATRIZ DE VARIÂNCIAS E COVARIÂNCIAS

A formulação hierárquica dos modelos mistos exige a especificação cuidadosa das estruturas de variâncias e covariâncias que compõem o modelo, uma vez que elas desempenham papel central na representação da variabilidade entre unidades experimentais e da dependência existente entre observações repetidas. Em termos gerais, a matriz de variância-covariância marginal, $\Sigma_i = \mathbf{Z}_i \mathbf{D} \mathbf{Z}_i^\top + \mathbf{R}_i$, resulta da combinação de duas componentes distintas: a matriz associada aos efeitos aleatórios, denotada por $\mathbf{D}$, e a matriz residual, denotada por $\mathbf{R}_i$ para o i -ésimo indivíduo. Cada uma dessas matrizes incorpora informações específicas sobre a estrutura estocástica do modelo e, portanto, necessita ser interpretada e especificada de acordo com a natureza do experimento.

A matriz $\mathbf{D}$ descreve a variabilidade entre os efeitos aleatórios, sendo definida como uma matriz simétrica e definida positiva que contém, em sua diagonal, as variâncias de cada componente aleatório e, fora dela, suas respectivas covariâncias. Como discutido por West, Welch e Galecki (2014), essa matriz é parametrizada por um conjunto reduzido de parâmetros θ_D , os quais determinam a estrutura de dependência entre os efeitos aleatórios e impõem restrições sobre suas variâncias e covariâncias. A escolha da estrutura adequada para $\mathbf{D}$ influencia diretamente a flexibilidade do modelo e a interpretação dos componentes associados às diferenças individuais.

De forma complementar, a matriz residual $\mathbf{R}_i$ representa a variabilidade intraunidade experimental e a dependência entre observações coletadas ao longo do tempo ou em condições similares. Assim como a matriz $\mathbf{D}$, $\mathbf{R}_i$ é simétrica e definida positiva, e sua especificação determina como a correlação e a heterogeneidade residual são modeladas. Conforme destacado por Pinheiro e Bates (2000), a estrutura residual pode assumir diferentes formas, desde modelos homocedásticos e independentes até estruturas mais complexas que incorporam autocorrelação temporal, heterocedasticidade ou ambas. A parametrização desses comportamentos é reunida no vetor θ_R , o qual define a maneira como variâncias e covariâncias residuais são organizadas.

A distinção entre as matrizes $\mathbf{D}$ e $\mathbf{R}_i$ é fundamental: enquanto a primeira captura a variabilidade entre unidades experimentais, a segunda especifica a dependência e a dispersão dentro de cada unidade. A correta especificação dessas estruturas é essencial para garantir estimativas consistentes e interpretações coerentes, especialmente em estudos com medidas repetidas e processos de natureza dinâmica, como ocorre em experimentos de produção de gás *in vitro*. As seções seguintes apresentam as principais estruturas utilizadas para cada uma dessas matrizes no contexto dos modelos não lineares mistos empregados neste trabalho.

1.2.1.1 ESTRUTURAS DE COVARIÂNCIA PARA A MATRIZ DE EFEITOS ALEATÓRIOS DO MODELO (**D**)

Em modelos de efeitos mistos aplicados a dados longitudinais, como aqueles oriundos de ensaios de produção de gás *in vitro*, é essencial especificar adequadamente a matriz de variância-covariância dos efeitos aleatórios, denotada por **D**. Essa matriz representa a variabilidade entre unidades experimentais (como réplicas ou frascos) e descreve tanto as variâncias individuais de cada efeito aleatório quanto as possíveis correlações entre eles (PINHEIRO; BATES, 2000).

De acordo com Pinheiro e Bates (2000), a matriz **D** deve ser especificada com base em pressupostos que reflitam adequadamente a heterogeneidade dos parâmetros aleatórios e a possível interdependência entre eles. A adoção de uma estrutura inadequada para **D** pode comprometer a qualidade do ajuste, levando a estimativas enviesadas dos efeitos fixos, erros-padrão incorretos e inferências estatísticas inválidas.

Nos modelos mistos, **D** pode assumir diferentes estruturas paramétricas, que são especificadas no ambiente estatístico R (R Core Team, 2025) por meio das classes da família `pdMat`, como `pdIdent`, `pdDiag` e `pdSymm` do pacote `nlme` (PINHEIRO; BATES; R Core Team, 2025). Essas estruturas oferecem flexibilidade na modelagem da variabilidade entre as unidades experimentais, permitindo diferentes formas de assumir ou não a independência e homogeneidade dos efeitos aleatórios.

Pinheiro e Bates (2000) destacam que a seleção da estrutura mais adequada para **D** pode ser feita por meio da comparação de modelos utilizando testes de razão de verossimilhança (quando apropriado) ou critérios de informação como AIC e BIC, sendo fundamental equilibrar parcimônia e qualidade do ajuste.

1.2.1.1.1 Estrutura identidade

A estrutura identidade (`pdIdent`) assume que todos os efeitos aleatórios possuem a mesma variância e são independentes entre si (PINHEIRO; BATES, 2000). Nesse caso, a matriz de variância-covariância dos efeitos aleatórios para o indivíduo i é expressa como:

$$\mathbf{D} = \sigma^2 \mathbf{I}_q, \quad (1.15)$$

em que σ^2 é a variância comum a todos os efeitos aleatórios e $\mathbf{I}_q$ representa a matriz identidade de ordem q , correspondente ao número de componentes de efeitos aleatórios especificados no modelo.

Essa estrutura é a mais simples e implica que não há correlação entre os efeitos aleatórios e que todos compartilham a mesma variabilidade. Sua utilização é recomendada quando não há indícios de heterogeneidade de variâncias ou correlação entre os parâmetros aleatórios.

1.2.1.1.2 Estrutura diagonal heterocedástica

A estrutura diagonal heterocedástica (**pdDiag**) permite que cada efeito aleatório possua uma variância própria, mas assume que não há correlação entre eles. Assim, a matriz de variância-covariância dos efeitos aleatórios é uma matriz diagonal, representada por:

$$\mathbf{D} = \begin{pmatrix} \sigma_1^2 & 0 & \cdots & 0 \\ 0 & \sigma_2^2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \sigma_q^2 \end{pmatrix}, \quad (1.16)$$

em que cada σ_j^2 representa a variância associada ao j -ésimo parâmetro aleatorizado, e q é o número de parâmetros de efeito aleatório no modelo.

De acordo com Pinheiro e Bates (2000), essa estrutura é adequada quando se suspeita da presença de heterocedasticidade entre os efeitos aleatórios, mas não se espera correlação entre eles.

1.2.1.1.3 Estrutura de simétrica não estruturada

A estrutura de simétrica não estruturada (**pdSymm**) representa a forma mais geral e flexível de especificação da matriz de variância-covariância dos efeitos aleatórios. Nessa configuração, nenhuma restrição é imposta além da simetria e da condição de ser uma matriz definida positiva. A matriz $\mathbf{D}$ para q parâmetros aleatorizados assume a forma:

$$\mathbf{D} = \begin{pmatrix} \sigma_1^2 & \sigma_{12} & \cdots & \sigma_{1q} \\ \sigma_{21} & \sigma_2^2 & \cdots & \sigma_{2q} \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_{q1} & \sigma_{q2} & \cdots & \sigma_q^2 \end{pmatrix}, \quad (1.17)$$

em que σ_j^2 são as variâncias dos efeitos aleatórios, e $\sigma_{jk} = \sigma_{kj}$ representam as covariâncias entre os efeitos aleatórios dos parâmetros j e k .

Essa estrutura permite capturar qualquer padrão de dependência entre os efeitos aleatórios, oferecendo máxima flexibilidade. No entanto, essa flexibilidade tem um custo, para q parâmetros aleatorizados, são necessários $\frac{q(q+1)}{2}$ parâmetros para especificar a matriz $\mathbf{D}$. Isso pode conduzir a problemas de sobreparametrização, especialmente em estudos com número limitado de unidades experimentais (PINHEIRO; BATES, 2000).

1.2.1.2 ESTRUTURAS DE COVARIÂNCIA PARA A MATRIZ RESIDUAL DO MODELO ($\mathbf{R}_i$)

A matriz residual $\mathbf{R}_i$ desempenha papel fundamental na modelagem de dados longitudinais, pois descreve a variabilidade e a dependência existente entre as observações coletadas dentro do mesmo indivíduo. Em modelos mistos, diferentes estruturas podem ser atribuídas a $\mathbf{R}_i$, permitindo acomodar desde situações com variância constante e ausência de correlação até cenários com padrões complexos de autocorrelação ou heterogeneidade ao longo do tempo. Como observado por West, Welch e Galecki (2014), a escolha apropriada da estrutura residual é determinante para representações adequadas da dependência temporal e para evitar vieses na estimação dos parâmetros do modelo.

1.2.1.2.1 Estrutura diagonal

A forma mais simples para a matriz residual consiste em assumir que os erros associados às observações de um mesmo indivíduo são independentes e possuem variância constante. Essa suposição resulta em uma matriz diagonal:

$$\mathbf{R}_i = \sigma^2 \mathbf{I}_{n_i}, \quad (1.18)$$

em que cada elemento da diagonal representa a variância residual e os termos fora da diagonal são nulos. Essa estrutura é adequada em contextos nos quais não há evidências de autocorrelação ou heterogeneidade ao longo do tempo. Entretanto, em dados longitudinais, essa condição é frequentemente irrealista, motivo pelo qual estruturas mais flexíveis costumam ser necessárias.

1.2.1.2.2 Simetria composta (CS)

A estrutura de simetria composta supõe que todas as observações de um mesmo indivíduo apresentam a mesma variância residual e compartilham um nível constante de correlação entre si. A matriz correspondente pode ser escrita como:

$$\mathbf{R}_i = \begin{bmatrix} \sigma^2 + \sigma_1 & \sigma_1 & \cdots & \sigma_1 \\ \sigma_1 & \sigma^2 + \sigma_1 & \cdots & \sigma_1 \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_1 & \sigma_1 & \cdots & \sigma^2 + \sigma_1 \end{bmatrix}, \quad (1.19)$$

em que σ^2 representa a variância específica de cada observação e σ_1 a covariância comum entre quaisquer dois resíduos. A hipótese de correlação constante é útil em experimentos com repetições sob condições idênticas, mas pode ser restritiva quando a dependência varia com o tempo entre medições (WEST; WELCH; GALECKI, 2014).

1.2.1.2.3 Estrutura de correlação exponencial (EXP)

Em estudos longitudinais com medições realizadas ao longo do tempo, é comum que observações mais próximas temporalmente apresentem maior semelhança do que aquelas mais distantes. Quando os tempos de avaliação não são igualmente espaçados, a modelagem da dependência residual deve considerar explicitamente a distância temporal entre as observações.

Nesse contexto, a estrutura de correlação exponencial (EXP), implementada pela função `corExp()` do pacote `nlme` (PINHEIRO; BATES; R Core Team, 2025), representa uma alternativa flexível para modelar a dependência temporal contínua entre os resíduos (PINHEIRO; BATES, 2000).

Nessa estrutura, a correlação entre dois resíduos associados aos tempos t' e t'' é descrita por:

$$\text{Corr}(\varepsilon_{it'}, \varepsilon_{it''}) = \exp\left(-\frac{|t' - t''|}{\rho}\right), \quad (1.20)$$

em que $\rho > 0$ representa o parâmetro de alcance (*range*), que controla a taxa de decaimento da correlação à medida que a distância temporal entre as observações aumenta.

Conforme apresentado por West, Welch e Galecki (2014), a estrutura exponencial pode ser definida como uma função de correlação decrescente em função da distância entre observações, em que a matriz de variância covariância residual associada à i -ésima unidade é expressa como:

$$\mathbf{R}_i = \sigma^2 \begin{bmatrix} 1 & \exp\left(-\frac{|t_{i1}-t_{i2}|}{\rho}\right) & \cdots & \exp\left(-\frac{|t_{i1}-t_{in_i}|}{\rho}\right) \\ \exp\left(-\frac{|t_{i2}-t_{i1}|}{\rho}\right) & 1 & \cdots & \exp\left(-\frac{|t_{i2}-t_{in_i}|}{\rho}\right) \\ \vdots & \vdots & \ddots & \vdots \\ \exp\left(-\frac{|t_{in_i}-t_{i1}|}{\rho}\right) & \exp\left(-\frac{|t_{in_i}-t_{i2}|}{\rho}\right) & \cdots & 1 \end{bmatrix}, \quad (1.21)$$

em que σ^2 representa a variância residual e n_i o número de observações da unidade experimental i .

Essa estrutura é particularmente adequada quando a dependência entre as observações diminui de forma contínua com o aumento da distância temporal, sem exigir espaçamento regular entre os tempos de medição. Dessa forma, a estrutura exponencial permite capturar de maneira realista a autocorrelação presente em dados longitudinais de produção de gases, contribuindo para maior estabilidade na estimação dos parâmetros e melhor ajuste dos modelos não lineares mistos.

1.2.1.2.4 Estrutura de correlação Gaussiana (GAUS)

Em estudos longitudinais, a dependência entre observações pode apresentar um padrão de decaimento mais suave nas proximidades e mais acentuado à medida que a distância temporal aumenta. Nesses casos, a estrutura de correlação Gaussiana (GAUS), implementada na função `corGaus()` do pacote `nlme` (PINHEIRO; BATES; R Core Team, 2025), constitui uma alternativa adequada para modelar a dependência residual contínua (PINHEIRO; BATES, 2000).

Essa estrutura pode ser compreendida como uma extensão da forma exponencial, na qual o decaimento da correlação ocorre de maneira mais rápida para distâncias maiores, sendo descrita por uma função quadrática da distância temporal (WEST; WELCH; GALECKI, 2014).

Nessa formulação, a correlação entre dois resíduos associados aos tempos t' e t'' é dada por:

$$\text{Corr}(\varepsilon_{it'}, \varepsilon_{it''}) = \exp\left(-\frac{(t' - t'')^2}{\rho^2}\right), \quad (1.22)$$

em que $\rho > 0$ representa o parâmetro de alcance (*range*), que controla a taxa de decaimento da correlação em função da distância temporal entre as observações.

A matriz de variância covariância residual associada à i -ésima unidade experimental pode ser expressa como:

$$\mathbf{R}_i = \sigma^2 \begin{bmatrix} 1 & \exp\left(-\frac{(t_{i1}-t_{i2})^2}{\rho^2}\right) & \cdots & \exp\left(-\frac{(t_{i1}-t_{in_i})^2}{\rho^2}\right) \\ \exp\left(-\frac{(t_{i2}-t_{i1})^2}{\rho^2}\right) & 1 & \cdots & \exp\left(-\frac{(t_{i2}-t_{in_i})^2}{\rho^2}\right) \\ \vdots & \vdots & \ddots & \vdots \\ \exp\left(-\frac{(t_{in_i}-t_{i1})^2}{\rho^2}\right) & \exp\left(-\frac{(t_{in_i}-t_{i2})^2}{\rho^2}\right) & \cdots & 1 \end{bmatrix}, \quad (1.23)$$

em que σ^2 representa a variância residual e n_i o número de observações da unidade experimental i .

Essa estrutura é particularmente útil quando a correlação entre observações decresce de forma mais acentuada à medida que a distância temporal aumenta, sendo adequada para descrever processos nos quais a dependência é mais forte entre observações muito próximas no tempo e se reduz rapidamente para intervalos maiores. Dessa forma, a estrutura Gaussiana permite capturar padrões de autocorrelação mais localizados, contribuindo para maior flexibilidade na modelagem da dependência residual em dados longitudinais.

1.2.1.3 ESTRUTURA DE VARIÂNCIA HETEROGÊNEA (ID)

Em modelos longitudinais, é comum que a variabilidade dos resíduos não seja constante ao longo dos níveis de um fator ou entre grupos experimentais, caracterizando a presença de heterocedasticidade. A suposição de variância residual constante pode, nesses casos, comprometer

a qualidade do ajuste e a validade das inferências, especialmente quando diferentes tratamentos ou unidades experimentais apresentam magnitudes distintas de variabilidade.

Para acomodar esse comportamento, foi considerada a utilização de estruturas de variância heterogênea por meio da função `varIdent()`, disponível no pacote `nlme` (PINHEIRO; BATES; R Core Team, 2025) assim como implementado em Pinheiro e Bates (2000). Essa estrutura permite modelar variâncias residuais distintas para diferentes níveis de um fator categórico, sem alterar a forma funcional da correlação entre as observações.

Nessa abordagem, a variância residual associada à observação j da unidade i pode ser expressa como:

$$\text{Var}(\varepsilon_{ij}) = \sigma^2 \cdot \delta_{g_{ij}}^2, \quad (1.24)$$

em que σ^2 representa a variância residual de referência e $\delta_{g_{ij}}$ é um fator de escala associado ao grupo g ao qual a observação pertence, sendo estimado a partir dos dados.

Assim, diferentes níveis de um fator (por exemplo, tratamentos ou combinações de fatores experimentais) podem apresentar distintas magnitudes de variância, permitindo maior flexibilidade na modelagem da dispersão dos resíduos. Sendo utilizada quando há evidência de heterogeneidade de variância entre grupos, contribuindo para melhor ajuste do modelo e maior consistência das inferências (PINHEIRO; BATES, 2000).

A incorporação da função `varIdent()` em conjunto com estruturas de correlação residual possibilita a modelagem simultânea da dependência temporal e da heterogeneidade de variância, resultando em uma representação mais realista da estrutura dos dados longitudinais.

1.3 MODELOS NÃO LINEARES MISTOS

Modelos não lineares mistos (MNLM) constituem uma extensão dos modelos lineares mistos, nos quais a relação entre a variável resposta e os parâmetros do modelo é definida por uma função não linear. Segundo Pinheiro e Bates (2000), essa classe de modelos é particularmente útil na análise de fenômenos com estrutura funcional intrinsecamente não linear, como curvas de crescimento, processos enzimáticos e produção de gás ao longo do tempo. Além disso, esses modelos permitem incorporar variabilidade entre indivíduos por meio de efeitos aleatórios, representando de forma mais realista os mecanismos subjacentes aos dados observados.

Para o i -ésimo indivíduo, com $i = 1, \dots, M$, e para cada observação $j = 1, \dots, n_i$, o modelo não linear misto assume a forma (LINDSTROM; BATES, 1990):

$$y_{ij} = f(\boldsymbol{\phi}_i, \mathbf{v}_{ij}) + \varepsilon_{ij}, \quad (1.25)$$

em que:

- y_{ij} representa a observação da variável resposta para o indivíduo i na j -ésima medição;
- $\mathbf{v}_{ij}$ é o vetor de covariáveis (por exemplo, tempo de incubação) associado à observação j do indivíduo i ;
- $f(\boldsymbol{\phi}_i, \mathbf{v}_{ij})$ é uma função não linear, real e diferenciável, aplicada ponto a ponto;
- $\boldsymbol{\phi}_i$ é o vetor de parâmetros do modelo não linear para a observação j do indivíduo i ;
- $\varepsilon_{ij} \sim N(0, \sigma^2)$ representa o erro intraindividual, assumido homocedástico e independente entre observações e indivíduos.

O vetor de parâmetros individuais $\boldsymbol{\phi}_i$ é especificado hierarquicamente como:

$$\boldsymbol{\phi}_i = \mathbf{A}_i \boldsymbol{\beta} + \mathbf{B}_i \mathbf{b}_i, \quad (1.26)$$

em que $\boldsymbol{\beta}$ representa o vetor de parâmetros fixos (populacionais), $\mathbf{b}_i \sim N(0, \sigma^2 \mathbf{D})$ é o vetor de efeitos aleatórios específicos do indivíduo i , e $\mathbf{A}_i$, $\mathbf{B}_i$ são matrizes de delineamento associadas aos efeitos fixos e aleatórios, respectivamente. Essa formulação, proposta por Lawrence e Lin (1989), permite descrever estruturas complexas de variabilidade entre e dentro de indivíduos, especialmente em contextos nos quais a resposta depende não linearmente dos parâmetros. Ao empilhar os dados de todos os indivíduos, o modelo pode ser reescrito como:

$$\mathbf{y}_i = \boldsymbol{\eta}_i(\boldsymbol{\phi}_i) + \boldsymbol{\varepsilon}_i, \quad (1.27)$$

em que:

$$\mathbf{y}_i = \begin{bmatrix} \mathbf{y}_{i1} \\ \mathbf{y}_{i2} \\ \vdots \\ \mathbf{y}_{in_i} \end{bmatrix}, \quad \boldsymbol{\varepsilon}_i = \begin{bmatrix} \varepsilon_{i1} \\ \varepsilon_{i2} \\ \vdots \\ \varepsilon_{in_i} \end{bmatrix}, \quad \boldsymbol{\eta}_i(\boldsymbol{\phi}_i) = \begin{bmatrix} f(\boldsymbol{\phi}_i, \mathbf{v}_{i1}) \\ f(\boldsymbol{\phi}_i, \mathbf{v}_{i2}) \\ \vdots \\ f(\boldsymbol{\phi}_i, \mathbf{v}_{in_i}) \end{bmatrix}, \quad \boldsymbol{\varepsilon}_i \sim N(\mathbf{0}, \sigma^2 \mathbf{R}_i).$$

Para simplificar a notação, suprime-se a dependência da média em relação aos vetores de preditores correspondentes às observações dentro de cada unidade experimental.

Em muitas aplicações, considera-se $\mathbf{A}_i$ igual à matriz identidade, o que implica independência entre as observações dentro da mesma unidade. Contudo, a inclusão explícita da matriz $\mathbf{A}_i$ permite especificar estruturas de covariância marginal não independentes, como aquelas que incorporam correlação serial ao longo do tempo, a exemplo de estruturas autorregressivas. Em geral, o número de parâmetros necessários para descrever essa matriz é relativamente pequeno.

Os M modelos individuais podem ser representados de forma conjunta, empilhando-se as observações e vetores de erros correspondentes, resultando em uma formulação matricial

global na qual a matriz de covariância assume estrutura bloco-diagonal, composta pelas matrizes $\mathbf{A}_i$ associadas a cada unidade experimental.

$$\mathbf{y} = \begin{bmatrix} \mathbf{y}_1 \\ \mathbf{y}_2 \\ \vdots \\ \mathbf{y}_M \end{bmatrix}, \quad \boldsymbol{\phi} = \begin{bmatrix} \phi_1 \\ \phi_2 \\ \vdots \\ \phi_M \end{bmatrix}, \quad \boldsymbol{\eta}(\boldsymbol{\phi}) = \begin{bmatrix} \eta_1(\phi_1) \\ \eta_2(\phi_2) \\ \vdots \\ \eta_M(\phi_M) \end{bmatrix}.$$

As matrizes de delineamento e covariância são definidas como:

$$\mathbf{R} = \text{diag}(\mathbf{R}_1, \mathbf{R}_2, \dots, \mathbf{R}_M), \quad \tilde{\mathbf{D}} = \text{diag}(\mathbf{D}, \dots, \mathbf{D}).$$

Portanto, o modelo hierárquico completo assume a forma:

$$\mathbf{y} \mid \mathbf{b} \sim \mathbf{N}(\boldsymbol{\eta}(\boldsymbol{\phi}), \sigma^2 \mathbf{R}), \quad (1.28)$$

$$\boldsymbol{\phi} = \mathbf{A}\boldsymbol{\beta} + \mathbf{B}\mathbf{b}, \quad (1.29)$$

$$\mathbf{b} \sim \mathbf{N}(\mathbf{0}, \sigma^2 \tilde{\mathbf{D}}), \quad (1.30)$$

sendo $\mathbf{B} = \text{diag}(\mathbf{B}_1, \mathbf{B}_2, \dots, \mathbf{B}_M)$, $\mathbf{b} = (\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_M)^\top$ e $\mathbf{A} = (\mathbf{A}_1, \mathbf{A}_2, \dots, \mathbf{A}_M)^\top$. A matriz de covariância marginal dos dados, que incorpora simultaneamente a variabilidade dos efeitos aleatórios e dos resíduos, é dada por:

$$\boldsymbol{\Sigma} = \mathbf{Z}\mathbf{D}\mathbf{Z}^\top + \mathbf{R}, \quad (1.31)$$

em que $\mathbf{Z}$ é a matriz estruturada de delineamento dos efeitos aleatórios para todos os indivíduos, $\mathbf{D}$ é a matriz de variância-covariância dos efeitos aleatórios, com possíveis parametrizações como `pdIdent`, `pdDiag` ou `pdSymm`, e $\mathbf{R}$ representa a estrutura de correlação residual intraindivíduo.

A formulação dos MNLM apresenta grande flexibilidade na modelagem de dados correlacionados com estrutura funcional complexa. Segundo Pinheiro e Bates (2000), essa abordagem é especialmente apropriada para experimentos biológicos com medidas repetidas, nos quais as respostas seguem padrões não lineares com significado fisiológico.

Em estudos de degradabilidade ruminal a resposta ao longo do tempo comumente assume uma forma sigmoide, refletindo três fases distintas: latência inicial, crescimento acelerado e estabilização assintótica. Tal padrão exige o uso de funções não lineares com parâmetros interpretáveis biologicamente, como o volume máximo de gás, a taxa de fermentação e o tempo de inflexão da curva.

Os MNLM permitem ajustar essas funções ao mesmo tempo em que capturam a variabilidade intra e interindivíduos por meio de efeitos aleatórios. Essa estrutura é robusta frente a

dados incompletos, desbalanceamentos e heterogeneidade de variâncias, favorecendo inferências mais realistas e eficientes.

1.3.1 ESTIMAÇÃO DE β E $\mathbf{b}$

A estimação dos parâmetros nos modelos não lineares mistos, sob a hipótese de linearidade da função de predição, pode ser tratada como um problema de mínimos quadrados generalizados (GLS). Considerando que os componentes de variância-covariância $\mathbf{D}$ e $\mathbf{R}$ são conhecidos e η é uma função linear de β e $\mathbf{b}$, então temos que $\eta(\phi_i) = \eta(\mathbf{A}_i\beta + \mathbf{B}_i\mathbf{b}_i) = \mathbf{X}_i\beta + \mathbf{Z}_i\mathbf{b}_i$, e define-se:

$$\mathbf{X} = \begin{bmatrix} \mathbf{X}_1^\top \\ \mathbf{X}_2^\top \\ \vdots \\ \mathbf{X}_n^\top \end{bmatrix} \quad \text{e} \quad \mathbf{Z} = \text{diag}(\mathbf{Z}_1, \mathbf{Z}_2, \dots, \mathbf{Z}_n).$$

Os estimadores usuais para os parâmetros β e $\mathbf{b}$ são dados pelo estimador GLS (LINDSTROM; BATES, 1990)

$$\hat{\beta}_{\text{lin}} = \hat{\beta}_{\text{lin}}(\theta) = (\mathbf{X}^\top \boldsymbol{\Sigma}^{-1} \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\Sigma}^{-1} \mathbf{y}, \quad (1.32)$$

$$\hat{\mathbf{b}}_{\text{lin}} = \hat{\mathbf{b}}_{\text{lin}}(\theta) = \tilde{\mathbf{D}} \mathbf{Z}^\top \boldsymbol{\Sigma}^{-1} (\mathbf{y} - \mathbf{X} \hat{\beta}_{\text{lin}}(\theta)), \quad (1.33)$$

sendo $\boldsymbol{\Sigma} = \mathbf{Z} \tilde{\mathbf{D}} \mathbf{Z}^\top + \mathbf{R}$ a matriz de variância covariância marginal do vetor de respostas, em que $\mathbf{D}$ representa a matriz de variância covariância dos efeitos aleatórios e $\mathbf{R}$ a matriz de variância covariância residual. O vetor paramétrico θ reúne os elementos que compõem as matrizes $\mathbf{D}$ e $\mathbf{R}$, sendo introduzido como uma forma conveniente de representar conjuntamente os parâmetros de covariância do modelo, sem a necessidade de explicitar a parametrização particular adotada para $\mathbf{R}$. Assim, θ é definido em função das estruturas de variabilidade especificadas para os efeitos aleatórios e para os resíduos.

As estimativas $\hat{\beta}_{\text{lin}}$ e $\hat{\mathbf{b}}_{\text{lin}}$ maximizam conjuntamente a função:

$$\mathbf{g}_{\text{lin}}(\beta, \mathbf{b} | \mathbf{y}) = -\frac{1}{2} \sigma^{-2} (\mathbf{y} - \mathbf{X}\beta - \mathbf{Z}\mathbf{b})^\top \mathbf{R}^{-1} (\mathbf{y} - \mathbf{X}\beta - \mathbf{Z}\mathbf{b}) - \frac{1}{2} \sigma^{-2} \mathbf{b}^\top \tilde{\mathbf{D}}^{-1} \mathbf{b}, \quad (1.34)$$

em que $\mathbf{g}_{\text{lin}}(\beta, \mathbf{b} | \mathbf{y})$ representa o critério objetivo penalizado obtido a partir da linearização local do modelo não linear misto. O critério da equação (1.34), para β fixo, representa o logaritmo da densidade a posteriori de $\mathbf{b}$ (desconsiderando constantes), e para $\mathbf{b}$ fixo, corresponde à log-verossimilhança de β (também até uma constante aditiva). Os dois termos presentes nessa equação compõem uma soma de quadrados e um termo quadrático em $\mathbf{b}$.

Ao transformar o termo quadrático $\mathbf{b}^\top \tilde{\mathbf{D}}^{-1} \mathbf{b}$ em uma forma equivalente a uma soma de quadrados, é possível tratar o problema de otimização inteiramente como um problema de mínimos quadrados (LINDSTROM; BATES, 1990). Esse procedimento facilita a transição para o caso não linear. Esse problema de mínimos quadrados ao aumentar o vetor de dados com "pseudo-dados" é dada da seguinte forma:

$$\tilde{\mathbf{y}} = \tilde{\mathbf{X}}\boldsymbol{\beta} + \tilde{\mathbf{Z}}\mathbf{b} + \tilde{\boldsymbol{\varepsilon}}, \quad (1.35)$$

sendo:

$$\tilde{\mathbf{y}} = \begin{bmatrix} \mathbf{R}^{-1/2} \mathbf{y} \\ \mathbf{0} \end{bmatrix}, \quad \tilde{\mathbf{X}} = \begin{bmatrix} \mathbf{R}^{-1/2} \mathbf{X} \\ \mathbf{0} \end{bmatrix}, \quad \tilde{\mathbf{Z}} = \begin{bmatrix} \mathbf{R}^{-1/2} \mathbf{Z} \\ \tilde{\mathbf{D}}^{-1/2} \end{bmatrix}, \quad \tilde{\boldsymbol{\varepsilon}} \sim \mathbf{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

Na formulação proposta por Lindstrom e Bates (1990), a penalização dos efeitos aleatórios é incorporada ao problema de estimação por meio da adição de pseudo-observações. Para isso, considera-se a decomposição de Cholesky da matriz de variância dos efeitos aleatórios, tal que:

$$\mathbf{D}^{-1/2} = \text{diag}(\mathbf{L}_1^\top, \mathbf{L}_2^\top, \dots, \mathbf{L}_M^\top),$$

em que cada $\mathbf{L}_i$ é uma matriz triangular superior tal que $\mathbf{D} = \mathbf{L}_i \mathbf{L}_i^\top$. De maneira análoga, aplica-se a decomposição de Cholesky à matriz de covariância dos resíduos intraindividuais, $\mathbf{A}_i^{-1/2}$, a fim de incorporar a estrutura residual ao problema de estimação por mínimos quadrados generalizados.

Dessa forma, no contexto dos modelos mistos não lineares, os estimadores de máxima verossimilhança $\boldsymbol{\beta}(\boldsymbol{\theta})$ e o estimador empírico $\mathbf{b}(\boldsymbol{\theta})$ são obtidos pela maximização da seguinte função objetivo:

$$\mathbf{g}_{\text{lin}}(\boldsymbol{\beta}, \mathbf{b} \mid \mathbf{y}) = -\frac{1}{2\sigma^2} \left[(\mathbf{y} - \boldsymbol{\eta}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\mathbf{b}))^\top \mathbf{R}^{-1} (\mathbf{y} - \boldsymbol{\eta}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\mathbf{b})) + \mathbf{b}^\top \tilde{\mathbf{D}}^{-1} \mathbf{b} \right]. \quad (1.36)$$

Para $\boldsymbol{\beta}$ fixo, $\mathbf{g}_{\text{lin}}(\boldsymbol{\beta}, \mathbf{b} \mid \mathbf{y})$ equivale ao logaritmo da densidade a posteriori de $\mathbf{b}$ (até uma constante aditiva), enquanto que para $\mathbf{b}$ fixo, corresponde à log-verossimilhança associada a $\boldsymbol{\beta}$.

Como no caso linear, esse critério pode ser reinterpretado como um problema de mínimos quadrados generalizados, com a criação de um vetor de dados aumentado (*pseudo-data*), conforme:

$$\tilde{\mathbf{y}} = \tilde{\boldsymbol{\eta}}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\mathbf{b}) + \tilde{\boldsymbol{\varepsilon}}, \quad (1.37)$$

em que:

$$\tilde{\mathbf{y}} = \begin{bmatrix} \mathbf{R}^{-1/2} \mathbf{y} \\ \mathbf{0} \end{bmatrix}, \quad \tilde{\boldsymbol{\varepsilon}} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I}), \quad \text{e} \quad \tilde{\boldsymbol{\eta}}(\mathbf{A}\boldsymbol{\beta}, \mathbf{B}\mathbf{b}) = \begin{bmatrix} \mathbf{R}^{-1/2} \boldsymbol{\eta}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\mathbf{b}) \\ \tilde{\mathbf{D}}^{-1/2} \mathbf{b} \end{bmatrix}.$$

Assim, a estimação conjunta de $\boldsymbol{\beta}$ e $\mathbf{b}$ no modelo misto não linear é reduzida à minimização da distância entre $\tilde{\mathbf{y}}$ e $\tilde{\boldsymbol{\eta}}(\boldsymbol{\beta}, \mathbf{b})$, isto é, um problema de mínimos quadrados não linear penalizado.

1.3.2 ESTIMAÇÃO DOS COMPONENTES DE VARIÂNCIA

1.3.2.1 MÁXIMA VEROSSIMILHANÇA (MV)

Nos modelos mistos não lineares, a estimação do vetor de parâmetros de dispersão $\boldsymbol{\theta} = (\sigma^2, \text{parâmetros de } \mathbf{D} \text{ e/ou } \mathbf{R})$, é realizada com base na maximização da função de verossimilhança marginal dos dados observados (LINDSTROM; BATES, 1990).

A densidade marginal de $\mathbf{y}$ é expressa como:

$$p(\mathbf{y}) = \int p(\mathbf{y} | \mathbf{b}) \cdot p(\mathbf{b}) d\mathbf{b}, \quad (1.38)$$

em que $p(\mathbf{y} | \mathbf{b})$ é a densidade condicional dos dados dado o vetor de efeitos aleatórios $\mathbf{b}$, e $p(\mathbf{b})$ é a densidade dos efeitos aleatórios.

No entanto, como a função não linear $\boldsymbol{\eta}(\cdot)$ depende de $\mathbf{b}$ de forma não linear, não existe uma solução analítica fechada para essa integral. Assim, segundo Lindstrom e Bates (1990), realiza-se uma aproximação local de $\mathbf{y} - \boldsymbol{\eta}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\mathbf{b})$ por uma expansão de Taylor de primeira ordem em torno do modo a posteriori dos efeitos aleatórios $\hat{\mathbf{b}}$.

A aproximação da diferença é dada por:

$$\mathbf{y} - \boldsymbol{\eta}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\mathbf{b}) \approx \mathbf{y} - \boldsymbol{\eta}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\hat{\mathbf{b}}) + \hat{\mathbf{Z}}(\mathbf{b} - \hat{\mathbf{b}}), \quad (1.39)$$

em que $\hat{\mathbf{Z}}$ é a matriz das derivadas parciais (Jacobiana) de $\boldsymbol{\eta}$ em relação a $\mathbf{b}$, avaliada em $\mathbf{b} = \hat{\mathbf{b}}$:

$$\hat{\mathbf{Z}} = \hat{\mathbf{Z}}(\boldsymbol{\theta}) = \left. \frac{\partial \boldsymbol{\eta}}{\partial \mathbf{b}^\top} \right|_{\hat{\boldsymbol{\beta}}, \hat{\mathbf{b}}}. \quad (1.40)$$

Dessa forma, a distribuição marginal aproximada de $\mathbf{y}$ é expressa por:

$$\mathbf{y} \sim \mathcal{N}(\boldsymbol{\eta}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\hat{\mathbf{b}}) - \hat{\mathbf{Z}}\hat{\mathbf{b}}, \sigma^2 \hat{\boldsymbol{\Sigma}}), \quad (1.41)$$

sendo:

$$\hat{\boldsymbol{\Sigma}} = \hat{\mathbf{Z}}\tilde{\mathbf{D}}\hat{\mathbf{Z}}^\top + \mathbf{R}. \quad (1.42)$$

A função de log-verossimilhança marginal correspondente é:

$$\ell(\boldsymbol{\beta}, \sigma, \boldsymbol{\theta} \mid \mathbf{y}) = -\frac{1}{2} \log |\sigma^2 \widehat{\boldsymbol{\Sigma}}| - \frac{1}{2\sigma^2} \left[\mathbf{y} - \boldsymbol{\eta}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\widehat{\mathbf{b}}) + \mathbf{Z}\widehat{\mathbf{b}} \right]^\top \widehat{\boldsymbol{\Sigma}}^{-1} \left[\mathbf{y} - \boldsymbol{\eta}(\mathbf{A}\boldsymbol{\beta} + \mathbf{B}\widehat{\mathbf{b}}) + \mathbf{Z}\widehat{\mathbf{b}} \right]. \quad (1.43)$$

A maximização de 1.43 é realizada numericamente em relação a $\boldsymbol{\theta}$, geralmente utilizando algoritmos iterativos como Newton-Raphson ou Fisher-Scoring. Após a obtenção de $\widehat{\boldsymbol{\theta}}$, os estimadores dos parâmetros fixos $\widehat{\boldsymbol{\beta}}$ são estimados e os preditores dos efeitos aleatórios $\widehat{\mathbf{b}}$ são obtidos por mínimos quadrados generalizados ponderados (LINDSTROM; BATES, 1990).

1.3.2.2 MÁXIMA VEROSSIMILHANÇA RESTRITA (MVR)

De acordo com Lindstrom e Bates (1990), o MVR é definido de forma análoga ao método de máxima verossimilhança (MV), diferenciando-se por incorporar uma correção que leva em consideração a perda de graus de liberdade associada à estimação dos efeitos fixos. A função de log-verossimilhança restrita é definida como:

$$\ell_R(\boldsymbol{\beta}, \sigma, \boldsymbol{\theta} \mid \mathbf{y}) = -\frac{1}{2} \log |\sigma^{-2} \widehat{\mathbf{P}}^\top \widehat{\boldsymbol{\Sigma}}^{-1} \widehat{\mathbf{P}}| + \ell(\boldsymbol{\beta}, \sigma, \boldsymbol{\theta} \mid \mathbf{y}), \quad (1.44)$$

em que, para a i -ésima unidade experimental,

$$\widehat{\mathbf{P}}_i = \widehat{\mathbf{P}}_i(\boldsymbol{\theta}) = \left. \frac{\partial \boldsymbol{\eta}_i}{\partial \boldsymbol{\beta}^\top} \right|_{\widehat{\boldsymbol{\beta}}, \widehat{\mathbf{b}}} = \left(\left. \frac{\partial \boldsymbol{\eta}_i}{\partial \boldsymbol{\phi}^\top} \right|_{\widehat{\boldsymbol{\beta}}, \widehat{\mathbf{b}}} \right) \mathbf{A}_i, \quad (1.45)$$

e o empilhamento das matrizes $\widehat{\mathbf{P}}_i$ resulta em

$$\widehat{\mathbf{P}} = \widehat{\mathbf{P}}(\boldsymbol{\theta}) = \begin{bmatrix} \widehat{\mathbf{P}}_1 \\ \widehat{\mathbf{P}}_2 \\ \vdots \\ \widehat{\mathbf{P}}_n \end{bmatrix} = \left. \frac{\partial \boldsymbol{\eta}}{\partial \boldsymbol{\beta}^\top} \right|_{\widehat{\boldsymbol{\beta}}, \widehat{\mathbf{b}}}. \quad (1.46)$$

O último termo, $\log |\sigma^2 \mathbf{X}^\top \boldsymbol{\Sigma}^{-1} \mathbf{X}|$, corresponde à penalização que ajusta a função de verossimilhança para refletir corretamente a perda de graus de liberdade devido à estimação dos efeitos fixos $\boldsymbol{\beta}$, sendo essa a principal diferença em relação à formulação da função de máxima verossimilhança (MV).

A estimação dos parâmetros de dispersão $\boldsymbol{\theta}$, dos componentes de variância e dos efeitos fixos $\boldsymbol{\beta}$ é realizada por meio da maximização da função de log-verossimilhança restrita (1.44). Esse processo é implementado via algoritmos numéricos iterativos, como Newton-Raphson ou Fisher scoring (LINDSTROM; BATES, 1990).

1.3.3 SELEÇÃO DE MODELOS E QUALIDADE DE AJUSTE

A etapa de seleção de modelos desempenha papel central na análise de dados ajustados por modelos não lineares mistos (MNLM). Como a função média é não linear e a estrutura de variabilidade envolve simultaneamente efeitos aleatórios e correlação residual, diferentes especificações podem produzir interpretações substancialmente distintas dos parâmetros biológicos e da dinâmica temporal capturada pelo modelo. Dessa forma, a definição de uma estratégia sistemática de comparação entre modelos é essencial para garantir inferências consistentes e, sobretudo, fisiologicamente plausíveis (PINHEIRO; BATES, 2000).

Em MNLM, a seleção de modelos envolve três dimensões principais:

- (i) a escolha da forma funcional da curva média;
- (ii) a definição da estrutura de variância-covariância dos efeitos aleatórios;
- (iii) a especificação da matriz residual, que determina o padrão de correlação intraindivíduo (DAVIDIAN; GILTINAN, 2003).

Cada uma dessas escolhas impacta diretamente a verossimilhança marginal, os diagnósticos de ajuste e a interpretação final dos parâmetros. Assim, a comparação entre modelos exige critérios formais que permitam ponderar simultaneamente ajuste e parcimônia.

Tradicionalmente, os critérios baseados em verossimilhança constituem a base para comparação entre modelos mistos, uma vez que a estimação por MV ou MVR fornece valores de log-verossimilhança que refletem tanto a adequação do modelo quanto a complexidade da estrutura de covariância utilizada. Entre os critérios mais empregados destacam-se:

- o teste de razão de verossimilhança para modelos aninhados;
- o critérios de informação de Akaike (AIC);
- o critério de Akaike corrigido (AIC_c), indicado para amostras pequenas;
- o critério de informação Bayesiano (BIC), que penaliza estruturas mais complexas.

Além desses critérios numéricos, a seleção de modelos em MNLM requer inspeção detalhada dos resíduos, verificação de autocorrelação remanescente, análise da variância condicional, e avaliação da capacidade preditiva marginal. Como enfatizam Pinheiro e Bates (2000), a análise gráfica desempenha papel tão importante quanto os critérios de verossimilhança, especialmente em situações nas quais a linearização local pode influenciar a forma da verossimilhança aproximada nos modelos não lineares.

Dessa forma, a seleção do modelo final neste estudo não se baseou em um único critério, mas na convergência de evidências provenientes de medidas de informação, testes formais, diagnósticos residuais, análise de autocorrelação e desempenho preditivo. Essa estratégia integrada,

em consonância com as recomendações de Medeiros et al. (2020), permite uma avaliação mais robusta da adequação estatística e da coerência biológica dos modelos ajustados aos dados de produção de gases *in vitro*.

1.3.3.1 TESTE DE RAZÃO DE VEROSSIMILHANÇA (TRV)

O teste de razão de verossimilhança (TRV) é uma ferramenta estatística utilizada para comparar modelos aninhados, ou seja, quando um modelo mais simples pode ser obtido a partir de um modelo mais complexo por meio da imposição de restrições nos parâmetros (PINHEIRO; BATES, 2000). A ideia central do TRV é quantificar a evidência fornecida pelos dados a favor do modelo completo em comparação ao modelo reduzido.

Sejam M_0 o modelo reduzido (hipótese nula H_0) e M_1 o modelo completo (hipótese alternativa H_1), e sejam ℓ_0 e ℓ_1 os respectivos valores da log-verossimilhança máxima sob cada modelo.

O cálculo da log-verossimilhança deve ser feito com base na máxima verossimilhança (MV) quando a comparação envolve modelos que diferem nos efeitos fixos. Por outro lado, quando os modelos diferem apenas na estrutura dos parâmetros de dispersão, como os efeitos aleatórios ou a matriz de variância-covariância dos resíduos, o uso da máxima verossimilhança restrita (MVR) é apropriado, pois essa abordagem fornece estimativas não viesadas para os componentes de variância, ajustando para os graus de liberdade dos efeitos fixos (PINHEIRO; BATES, 2000).

O estatístico de teste é definido por:

$$\Lambda = -2\log\left(\frac{\ell_1}{\ell_0}\right) = -2(\ell_0 - \ell_1), \quad (1.47)$$

em que Λ segue, sob H_0 , uma distribuição assintótica qui-quadrado com g graus de liberdade, sendo g o número de parâmetros adicionais no modelo completo em relação ao modelo reduzido (PINHEIRO; BATES, 2000). O valor de Λ é então comparado ao quantil da distribuição χ^2_g para o nível de significância α adotado. Se $\Lambda > \chi^2_{g,\alpha}$, rejeita-se a hipótese nula em favor do modelo mais complexo.

A comparação entre os modelos deve respeitar uma ordem hierárquica crescente de complexidade, na qual modelos reduzidos são comparados sequencialmente a modelos mais gerais. Essa estratégia, amplamente recomendada na literatura sobre análise de dados longitudinais, assegura a validade assintótica do teste e evita comparações entre modelos não aninhados ou estruturalmente incompatíveis (ANDRADE; SINGER, 1986).

No contexto de modelos mistos não lineares, essa hierarquização implica comparar, por exemplo, modelos com menor número de termos de efeitos aleatórios ou estruturas de covariância mais simples com modelos que incorporam maior flexibilidade estrutural. Tal procedimento é

explicitamente adotado por Pinheiro e Bates (2000) e operacionalizado em estudos aplicados, como no trabalho de Janeiro et al. (2022), em que a seleção do modelo procede da forma reduzida para a completa, respeitando a estrutura de aninhamento entre as especificações.

Cabe destacar que, quando o teste de razão de verossimilhança envolve a nulidade de componentes de variância, como em hipóteses do tipo $H_0 : \sigma^2 = 0$ versus $H_1 : \sigma^2 > 0$, a condição regular de que os parâmetros pertençam ao interior do espaço paramétrico não é satisfeita, uma vez que variâncias são restritas a valores não negativos. Nessas situações, conforme demonstrado por Self e Liang (1987), a estatística do teste não converge para uma distribuição qui-quadrado padrão, mas para uma combinação de distribuições qui-quadrado. Como consequência, o uso da distribuição qui-quadrado usual pode tornar o teste conservador, devendo-se interpretar os resultados com cautela quando a comparação envolver exclusivamente componentes de variância.

1.3.3.2 CRITÉRIOS DE INFORMAÇÃO: AIC, AIC_c E BIC

Além da razão de verossimilhança, a escolha entre modelos concorrentes pode ser feita com base em critérios de informação, que penalizam a complexidade do modelo para evitar sobreajuste (*overfitting*) (PINHEIRO; BATES, 2000). Dentre os critérios mais amplamente utilizados estão o critério de informação de Akaike (AIC), sua versão corrigida (AIC_c) e o critério de informação Bayesiano (BIC).

Esses critérios são derivados a partir da log-verossimilhança do modelo ajustado e incluem um termo penalizador que leva em consideração o número de parâmetros estimados. A ideia subjacente é encontrar um equilíbrio entre o bom ajuste aos dados (via log-verossimilhança) e a parcimônia do modelo (via penalização pela complexidade) (PINHEIRO; BATES, 2000).

O critério proposto por Akaike (1974) é definido por:

$$AIC = -2 \cdot \ell(\hat{\beta}, \hat{\theta}, \hat{\sigma}) + 2k, \quad (1.48)$$

sendo $\ell(\hat{\beta}, \hat{\theta}, \hat{\sigma})$ o valor da log-verossimilhança maximizada e k o número de parâmetros estimados no modelo. O AIC busca maximizar a qualidade preditiva do modelo, favorecendo aqueles com melhor ajuste, mas penalizando aqueles com muitos parâmetros, evitando a superparametrização. O modelo com menor valor de AIC é considerado preferível.

Quando o tamanho amostral é pequeno ou moderado, especialmente quando a razão $n/k < 40$, recomenda-se a utilização da versão corrigida do AIC, denominada AIC_c , proposta por Burnham e Anderson (2002). O AIC_c é calculado como:

$$AIC_c = AIC + \frac{2k(k+1)}{n-k-1}, \quad (1.49)$$

em que n é o número total de observações e k o número de parâmetros. A penalização adicional corrige o viés do AIC quando n não é suficientemente grande em relação ao número de parâmetros. Quando n tende ao infinito, o Critério de Informação de Akaike Corrigido (AICc) AIC_c converge para o AIC.

O BIC, definido por Schwarz (1978), adota uma penalização mais severa do que o AIC:

$$BIC = -2 \cdot \ell(\hat{\beta}, \hat{\theta}, \hat{\sigma}) + k \cdot \log(n), \quad (1.50)$$

em que n é o número de observações e k o número de parâmetros. O BIC é derivado de uma perspectiva bayesiana e tende a favorecer modelos mais parcimoniosos, sendo especialmente útil quando o objetivo é identificar o modelo verdadeiro dentro de um conjunto de modelos candidatos.

Conforme discutido por Burnham e Anderson (2002), o AIC (ou AIC_c) é mais adequado quando o objetivo é selecionar o modelo com melhor capacidade preditiva, enquanto o BIC é mais conservador, priorizando modelos mais simples, sobretudo quando o tamanho amostral é grande.

1.3.3.3 INTERVALOS DE CONFIANÇA NOS MODELOS NÃO LINEARES MISTOS

A construção de intervalos de confiança em modelos não lineares mistos segue princípios análogos aos utilizados nos modelos lineares mistos, embora as aproximações envolvidas sejam influenciadas pela linearização local do modelo em torno dos estimadores obtidos via máxima verossimilhança (MV) ou máxima verossimilhança restrita (MVR). Segundo Pinheiro e Bates (2000), tais aproximações permitem derivar intervalos para os efeitos fixos, para o desvio-padrão residual e para os componentes da matriz de variância-covariância dos efeitos aleatórios, ainda que cada classe de parâmetros exija tratamento diferenciado devido às restrições impostas pela estrutura hierárquica do modelo.

O vetor de efeitos fixos β é estimado a partir da linearização da função não linear $\eta(\phi)$ em torno dos valores atuais dos parâmetros. A variabilidade associada ao estimador $\hat{\beta}$ decorre da matriz de informação observada, avaliada nos estimadores MV ou MVR. Assim, um intervalo de confiança aproximado de nível $1 - \alpha$ para o componente β_j pode ser expresso como:

$$\hat{\beta}_j \pm t_{df_j}(1 - \alpha/2) \hat{\sigma}_R \sqrt{[\mathbf{R}_{00}^{-1} \mathbf{R}_{00}^{-T}]_{jj}}, \quad (1.51)$$

em que $\hat{\sigma}_R$ denota a estimativa de σ por MVR, $t_{df_j}(1 - \alpha/2)$ é o quantil da distribuição t associado aos graus de liberdade determinados pela estrutura do modelo e $\mathbf{R}_{00}$ denota o bloco triangular superior associado aos efeitos fixos na fatoração da matriz de informação. A matriz de variância aproximada de $\hat{\beta}$ é então obtida a partir do produto $\mathbf{R}_{00}^{-1} \mathbf{R}_{00}^{-T}$, garantindo estabilidade numérica na inversão e na obtenção dos erros padrão, conforme proposto por Pinheiro e Bates

(2000). A precisão desse intervalo depende da qualidade da aproximação local utilizada na estimação do modelo não linear misto.

O parâmetro σ , que rege a variabilidade intraindivíduo associada ao termo de erro $\varepsilon_i \sim N(0, \sigma^2 \mathbf{R}_i)$, possui distribuição assintótica aproximadamente log-normal sob MV ou MVR. A partir disso, Pinheiro e Bates (2000) deduzem um intervalo de confiança multiplicativo para σ , obtido a partir do elemento associado a esse parâmetro na matriz de informação empírica inversa:

$$\left[\hat{\sigma} \exp\left(-z(1 - \alpha/2)\sqrt{[\mathcal{I}^{-1}]_{\sigma\sigma}}\right), \hat{\sigma} \exp\left(z(1 - \alpha/2)\sqrt{[\mathcal{I}^{-1}]_{\sigma\sigma}}\right) \right], \quad (1.52)$$

em que $[\mathcal{I}^{-1}]_{\sigma\sigma}$ corresponde ao elemento da matriz de informação empírica associado a σ . Esta forma multiplicativa assegura que o intervalo permaneça na escala positiva, preservando a interpretação do desvio-padrão residual.

A matriz $\mathbf{D}$, que representa a variabilidade interindividual por meio da relação $\mathbf{b}_i \sim N(0, \sigma^2 \mathbf{D})$, está sujeita à restrição de ser simétrica e definida positiva. Isso impede que intervalos de confiança sejam obtidos diretamente na escala dos parâmetros irrestritos utilizados durante a otimização numérica, especialmente quando estruturas como `pdSymm` são utilizadas.

Para contornar esse problema, Pinheiro e Bates (2000) propõem a *parametrização natural*, que representa os desvios-padrão em escala logarítmica e as correlações por meio do *generalized logit*,

$$\eta_j = \log\left(\frac{1 + \rho_j}{1 - \rho_j}\right), \quad (1.53)$$

permitindo que cada componente do vetor natural $\boldsymbol{\eta}$ varie livremente na reta real. Após obter o intervalo de confiança nessa escala irrestrita, aplica-se a transformação inversa, resultando em intervalos válidos na escala original.

Assim, para um desvio-padrão associado a um elemento diagonal de $\mathbf{D}$, o intervalo aproximado de confiança é dado por:

$$\left(\exp\left[\hat{\eta}_j - z(1 - \alpha/2)\sqrt{[\mathcal{I}^{-1}]_{jj}}\right], \exp\left[\hat{\eta}_j + z(1 - \alpha/2)\sqrt{[\mathcal{I}^{-1}]_{jj}}\right] \right), \quad (1.54)$$

enquanto para uma correlação ρ_j representada por η_j , o intervalo correspondente na escala original é obtido pela inversão do logito generalizado:

$$\left(\frac{\exp\left[\hat{\eta}_j - z(1 - \alpha/2)\sqrt{[\mathcal{I}^{-1}]_{jj}}\right] - 1}{\exp\left[\hat{\eta}_j - z(1 - \alpha/2)\sqrt{[\mathcal{I}^{-1}]_{jj}}\right] + 1}, \frac{\exp\left[\hat{\eta}_j + z(1 - \alpha/2)\sqrt{[\mathcal{I}^{-1}]_{jj}}\right] - 1}{\exp\left[\hat{\eta}_j + z(1 - \alpha/2)\sqrt{[\mathcal{I}^{-1}]_{jj}}\right] + 1} \right). \quad (1.55)$$

Esta abordagem assegura que os limites estimados respeitem a estrutura de matriz definida positiva, além de possibilitar interpretações coerentes para os componentes de variância e correlação entre os efeitos aleatórios do modelo.

1.3.3.4 ANÁLISE DE RESÍDUOS

A análise de resíduos é uma etapa essencial na avaliação da qualidade do ajuste de modelos mistos, permitindo investigar possíveis violações dos pressupostos do modelo, identificar *outliers*, heterocedasticidade, autocorrelação e avaliar a adequação da estrutura funcional escolhida.

Segundo Pinheiro e Bates (2000), os resíduos podem ser definidos como as diferenças entre os valores observados e os valores ajustados pelo modelo, ou seja:

$$\varepsilon_{ij} = y_{ij} - \hat{y}_{ij}, \quad (1.56)$$

em que y_{ij} representa a j -ésima observação do indivíduo i e $\hat{y}_{ij}$ é o valor predito pelo modelo ajustado.

Os resíduos studentizados condicionais são definidos como a razão entre o resíduo individual e o desvio padrão estimado associado a esse resíduo, levando em consideração os efeitos fixos e aleatórios. Segundo (DEMIDENKO, 2013), essa medida ajusta o resíduo pelo grau de incerteza associado à sua predição, permitindo detectar observações influentes ou discrepantes.

A definição matemática é dada por:

$$r_i = \frac{y_i - \hat{y}_i}{\sqrt{\text{Var}(y_i - \hat{y}_i)}} \quad (1.57)$$

em que y_i é a observação do indivíduo i , $\hat{y}_i$ é o valor ajustado pela predição condicional isto é, com base em $\hat{\beta}$ e $\hat{\mathbf{b}}$ e $\text{Var}(y_i - \hat{y}_i)$ é a variância condicional estimada do erro.

Essa forma de resíduo é útil para fins de diagnóstico, pois considera tanto a estrutura de correlação quanto a heterocedasticidade dos dados. A interpretação dos resíduos studentizados é análoga à dos modelos lineares clássicos: valores absolutos superiores a 2 indicam potenciais *outliers* ou pontos mal ajustados, embora o critério dependa do número de observações e do contexto.

Em estudos longitudinais, como os ensaios de produção de gás para avaliação da degradabilidade ruminal, as observações realizadas ao longo do tempo em uma mesma unidade experimental tendem a ser correlacionadas. Essa dependência temporal entre as medidas sucessivas, quando não adequadamente tratada, manifesta-se por meio da autocorrelação dos

resíduos, o que viola a suposição de independência usual em modelos lineares clássicos e pode comprometer a validade das inferências estatísticas.

De acordo com Pinheiro e Bates (2000), a presença de autocorrelação nos resíduos indica que o modelo não está descrevendo integralmente a estrutura de dependência dos dados. Isso pode ocorrer, por exemplo, quando a matriz de covariância residual $\mathbf{R}_i$ é especificada de forma simplificada, assumindo erros independentes com variância constante (estrutura identidade), enquanto na prática existe dependência temporal entre as observações dentro de cada unidade experimental.

Nos modelos mistos, a estrutura de correlação dos resíduos pode ser representada explicitamente por meio da especificação da matriz $\mathbf{R}_i$, utilizando funções como `corCompSymm` (simetria composta), `corGaus` (Gaussiana) ou outras estruturas implementadas no pacote `nlme` (PINHEIRO; BATES; R Core Team, 2025) do software estatístico *R* (??). Além disso, a variabilidade dos resíduos pode ser modelada por meio de funções de variância, como `varIdent`, permitindo acomodar heterocedasticidade entre grupos. Essa especificação conjunta possibilita descrever adequadamente tanto a autocorrelação quanto a heterogeneidade da variância dos erros intraindividuais, contribuindo para a melhoria do ajuste e da robustez das inferências.

A avaliação da autocorrelação dos resíduos foi realizada por meio da função de autocorrelação dos resíduos condicionais (ACF), conforme abordagem sugerida por Pinheiro e Bates (2000). Essa análise consiste em examinar a presença de dependência serial não capturada pelo modelo ajustado, sendo fundamental no contexto de dados longitudinais com medidas repetidas ao longo do tempo.

Brockwell e Davis (2002) definem a função de autocorrelação para uma série de resíduos $\{\varepsilon_t\}_{t=1}^T$ como:

$$\rho(h) = \frac{\sum_{t=1}^{T-h} (\varepsilon_t - \bar{\varepsilon})(\varepsilon_{t+h} - \bar{\varepsilon})}{\sum_{t=1}^T (\varepsilon_t - \bar{\varepsilon})^2}, \quad (1.58)$$

em que $\rho(h)$ representa o coeficiente de autocorrelação para a defasagem h , ε_t é o resíduo no tempo t , e $\bar{\varepsilon}$ é a média dos resíduos. Essa função mede a correlação linear entre pares de resíduos separados por uma defasagem h .

Conforme destacado por Brockwell e Davis (2002), a função de autocorrelação é amplamente utilizada em séries temporais para detectar padrões de dependência ao longo do tempo. No contexto dos modelos mistos, a aplicação dessa função aos resíduos condicionais permite verificar se a suposição de independência residual é adequada. Caso contrário, indica a necessidade de especificação de uma estrutura de correlação apropriada para a matriz de variância-covariância dos resíduos $\mathbf{R}$.

Assim, a utilização da função ACF, baseada na definição clássica da teoria de séries temporais proposta por Brockwell e Davis (2002), concentra-se na interpretação gráfica dos

resíduos para avaliar a adequação dos modelos ajustados.

1.3.3.5 AVALIAÇÃO DA CAPACIDADE PREDITIVA EM MODELOS NÃO LINEARES

No contexto da modelagem não linear, a avaliação da capacidade preditiva dos modelos constitui um aspecto fundamental, especialmente quando o objetivo é extrapolar os resultados para condições não observadas. Diferentemente das medidas baseadas exclusivamente no ajuste aos dados utilizados na estimação, a avaliação preditiva requer a verificação do desempenho do modelo em dados independentes, não utilizados no processo de ajuste.

De acordo com Hastie, Tibshirani e Friedman (2009), a distinção entre erro de treinamento e erro de generalização é central na avaliação de modelos estatísticos. Modelos que apresentam excelente ajuste aos dados observados podem não manter o mesmo desempenho quando aplicados a novos dados, caracterizando o fenômeno conhecido como sobreajuste (*overfitting*).

Nesse contexto, a capacidade preditiva de um modelo está associada à sua habilidade de generalizar padrões além da amostra utilizada na estimação. Idealmente, essa avaliação deve ser realizada com base em dados independentes. Na ausência desses, procedimentos como validação cruzada ou a divisão da amostra em conjuntos de treinamento e teste podem ser utilizados como aproximações para estimar o erro de generalização.

Na comparação entre funções não lineares que descrevem um mesmo fenômeno biológico, a literatura recomenda que a seleção de modelos não se baseie exclusivamente em critérios de verossimilhança, mas também na avaliação da adequação do ajuste aos dados e na interpretabilidade dos parâmetros estimados. Em modelagem aplicada às ciências agrárias e biológicas, especialmente em estudos envolvendo curvas de crescimento, degradação ou resposta acumulativa, é comum a utilização de diferentes critérios complementares para a escolha do modelo mais adequado (ARCHONTOULIS; MIGUEZ, 2015).

Segundo Archontoulis e Miguez (2015), a avaliação do desempenho de funções não lineares alternativas deve considerar não apenas a qualidade do ajuste global, mas também aspectos relacionados à coerência biológica dos parâmetros e ao comportamento dos resíduos.

A avaliação da capacidade preditiva em modelos não lineares pode também ser conduzida por meio da comparação entre valores observados e estimados, utilizando representações gráficas que permitem analisar a concordância entre essas quantidades. Dentre essas abordagens, destaca-se o gráfico de valores preditos versus observados, amplamente empregado como ferramenta diagnóstica em estudos aplicados.

Nesse tipo de representação, a proximidade dos pontos em relação à linha de identidade indica boa concordância entre valores observados e preditos, enquanto padrões sistemáticos de desvio podem sinalizar inadequações no modelo. Essa abordagem, embora não substitua procedimentos formais de validação externa, constitui uma forma complementar de avaliação

preditiva, especialmente quando dados independentes não estão disponíveis (HASTIE; TIBSHIRANI; FRIEDMAN, 2009).

Além das abordagens gráficas, a avaliação da capacidade preditiva de modelos não lineares pode ser conduzida por meio de estatísticas baseadas na discrepância entre valores observados e estimados. Dentre essas medidas, destacam-se a raiz do erro quadrático médio (RMSE), o erro absoluto médio (MAE) e o erro percentual absoluto médio (MAPE), amplamente utilizados para quantificar a magnitude dos desvios preditivos (ARCHONTOULIS; MIGUEZ, 2015).

O RMSE, em particular, é uma das métricas mais empregadas na comparação entre modelos, por penalizar de forma mais intensa erros de maior magnitude, sendo sensível à presença de desvios extremos. Já o MAE fornece uma medida de erro médio menos influenciada por valores discrepantes, enquanto o MAPE permite interpretar os erros em termos percentuais, facilitando a comparação entre diferentes escalas de resposta (apresentados na Tabela 1).

Além das medidas de erro, a avaliação da concordância entre valores observados e preditos pode ser realizada por meio do coeficiente de correlação de Pearson e do coeficiente de concordância de Lin (CCC) (Tabela 1). Enquanto a correlação de Pearson avalia a associação linear entre as variáveis, o CCC integra simultaneamente precisão e exatidão, sendo recomendado quando o interesse está na concordância entre valores observados e estimados (LAWRENCE; LIN, 1989).

Em modelos mistos não lineares, a avaliação preditiva pode ainda considerar diferentes níveis de predição, distinguindo-se entre predições marginais (baseadas apenas nos efeitos fixos) e predições condicionais (que incorporam também os efeitos aleatórios). Essa distinção é particularmente relevante em dados longitudinais, pois permite avaliar tanto o comportamento médio da população quanto a capacidade do modelo em descrever trajetórias específicas das unidades experimentais (PINHEIRO; BATES, 2000).

Dessa forma, a avaliação da capacidade preditiva em modelos não lineares pode ser conduzida de maneira integrada, combinando métricas de erro, medidas de concordância e representações gráficas, como o gráfico de valores observados versus preditos. Essa abordagem conjunta possibilita uma análise mais abrangente do desempenho dos modelos, contemplando tanto a qualidade do ajuste quanto a capacidade de representação dos dados observados.

Tabela 1 – Métricas utilizadas para avaliação do desempenho preditivo.

Métrica	Expressão matemática
Raiz do Erro Quadrático Médio (RMSE)	$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2}$
Erro Absoluto Médio (MAE)	$MAE = \frac{1}{n} \sum_{i=1}^n y_i - \hat{y}_i $
Erro Percentual Absoluto Médio (MAPE)	$MAPE = 100 \times \frac{1}{n} \sum_{i=1}^n \left \frac{y_i - \hat{y}_i}{y_i} \right $
Correlação de Pearson (Corr)	$Corr = \frac{\sum_{i=1}^n (y_i - \bar{y})(\hat{y}_i - \bar{\hat{y}})}{\sqrt{\sum_{i=1}^n (y_i - \bar{y})^2 \sum_{i=1}^n (\hat{y}_i - \bar{\hat{y}})^2}}$
Coeficiente de Concordância (CCC)	$CCC = \frac{2 \text{Cov}(y, \hat{y})}{\text{Var}(y) + \text{Var}(\hat{y}) + (\bar{y} - \bar{\hat{y}})^2}$

y_i representa o valor observado; $\hat{y}_i$ o valor predito; $\bar{y}$ e $\bar{\hat{y}}$ as respectivas médias; n o número total de observações; $\text{Cov}(y, \hat{y})$ a covariância entre valores observados e preditos; $\text{Var}(y)$ e $\text{Var}(\hat{y})$ as variâncias correspondentes.

MATERIAIS E MÉTODOS

2.1 MATERIAIS

Os dados utilizados nesta pesquisa são provenientes de um experimento conduzido no Laboratório de Zootecnia da Universidade Estadual de Maringá (UEM), empregando a técnica de degradabilidade ruminal *in vitro*. O ensaio foi realizado com fluido ruminal obtido de uma vaca Holandesa canulada, não lactante, com peso corporal médio de 545 kg, mantida em baia individual, conforme os critérios éticos aprovados pelo Comitê de Ética no Uso de Animais da UEM (protocolo CEUA/UEM nº 8208090218). Para fins ilustrativos, a Figura 1 apresenta um exemplo de animal canulado, destacando a etapa de coleta do fluido ruminal utilizada como inóculo biológico em ensaios *in vitro*.

Figura 1 – Animal canulado utilizado para a coleta de fluido ruminal, empregado como inóculo biológico em ensaios de degradabilidade ruminal *in vitro*.



Fonte: adaptado de Tomich et al. (2006).

Nota: a imagem refere-se exclusivamente à etapa de obtenção do inóculo ruminal, não correspondendo à incubação *in vitro*, a qual foi realizada em frascos sob condições controladas de temperatura, agitação e anaerobiose.

O fluido ruminal foi coletado diariamente após a primeira refeição, filtrado em gaze,

saturado com gás carbônico (CO₂) e misturado a uma solução de saliva artificial na proporção 20:80 (v/v), sendo mantido à temperatura constante de 39 °C. Em seguida, o inóculo ruminal tamponado foi utilizado para incubação das amostras de forragem em um sistema automatizado de produção de gases.

Foram avaliadas 16 amostras de forragens oriundas de diferentes regiões da América do Sul (Brasil, Argentina, Colômbia e Peru), incluindo silagens de milho, cana-de-açúcar e gramíneas tropicais, bem como fenos e pré-secados de alfafa, azevém, aveia e quicuí. As amostras foram secas em estufa de ventilação forçada a 55 °C por 72 horas no local de coleta, acondicionadas em sacos plásticos e enviadas à UEM para processamento e análises laboratoriais.

Para a condução do experimento *in vitro*, alíquotas de 500 mg de cada forragem foram acondicionadas em frascos de vidro, com três repetições por alimento. Cada frasco recebeu 60 mL de solução ruminal tamponada e foi mantido em banho-maria a 39 °C, sob agitação constante e ambiente anaeróbico. A produção de gás foi registrada automaticamente a cada 12 minutos durante 72 horas de incubação.

A produção cumulativa foi corrigida pela subtração do volume de gás proveniente de frascos controle (sem substrato), permitindo isolar a fração resultante exclusivamente da fermentação microbiana dos alimentos.

Os dados resultantes consistem em curvas individuais de produção cumulativa de gás ao longo do tempo para cada amostra de forragem. Cada unidade experimental gerou múltiplas observações sequenciais, configurando uma estrutura longitudinal com medidas repetidas, caracterizada por dependência intraunidade e possível autocorrelação serial entre observações próximas no tempo.

2.2 MÉTODOS

As análises estatísticas foram realizadas no ambiente computacional estatístico *R* (versão 4.3.1), utilizando-se principalmente o pacote *nlme* (PINHEIRO; BATES; R Core Team, 2025). O ajuste dos modelos não lineares de efeitos mistos foi conduzido por meio da função *nlme()*, conforme a formulação apresentada por Pinheiro e Bates (2000), a qual permite a incorporação simultânea de efeitos fixos, efeitos aleatórios e diferentes estruturas de variância-covariância para dados longitudinais.

Previamente ao ajuste dos modelos, os dados foram submetidos a procedimentos de tratamento e controle de qualidade, incluindo o truncamento de curvas inconsistentes, exclusão de observações atípicas e definição do esquema de tempos de incubação adotado, conforme descrito na Seção 1.1.1. Essas etapas visaram garantir maior consistência biológica e estabilidade numérica no processo de estimação.

Considerando que os modelos cinéticos e suas respectivas parametrizações já foram

apresentados no capítulo anterior, esta etapa concentrou-se na especificação da estrutura hierárquica dos modelos não lineares mistos, bem como na avaliação da necessidade e da forma de inclusão dos termos de efeitos aleatórios e da estrutura de dependência residual. Essa abordagem segue as recomendações de Lindstrom e Bates (1990) e Pinheiro e Bates (2000), segundo as quais a modelagem adequada da variabilidade interindividual é fundamental para a obtenção de estimativas consistentes e biologicamente interpretáveis em estudos com medidas repetidas.

De forma geral, o modelo adotado para descrever a produção de gás na i -ésima réplica, na j -ésima amostra, pode ser representado por

$$y_{ij} = f(\phi_i, t_{ij}) + \varepsilon_{ij}, \quad (2.1)$$

em que $f(\cdot)$ representa uma função não linear dependente do vetor de parâmetros individuais ϕ_i e do tempo de incubação t_{ij} , e ε_{ij} denota o erro aleatório associado à observação y_{ij} , assumido como $\varepsilon_{ij} \sim N(0, \sigma^2)$.

O vetor de parâmetros individuais ϕ_i é especificado hierarquicamente segundo a estrutura clássica dos modelos mistos não lineares:

$$\phi_i = \mathbf{A}_i \boldsymbol{\beta} + \mathbf{B}_i \mathbf{b}_i, \quad \mathbf{b}_i \sim N(\mathbf{0}, \mathbf{D}), \quad (2.2)$$

em que $\boldsymbol{\beta}$ representa o vetor de parâmetros fixos populacionais, $\mathbf{b}_i$ o vetor de efeitos aleatórios específicos da unidade experimental i , e $\mathbf{A}_i$ e $\mathbf{B}_i$ são as matrizes de delineamento associadas aos efeitos fixos e aleatórios, respectivamente. A matriz $\mathbf{D}$ descreve a estrutura de variância-covariância dos efeitos aleatórios entre unidades experimentais, enquanto a matriz $\mathbf{R}$ representa a estrutura de variância-covariância residual associada às observações intraindivíduo.

A partir dessa formulação geral, procedeu-se à avaliação da necessidade de inclusão de termos de efeitos aleatórios associados aos parâmetros das funções cinéticas. Ressalta-se que os parâmetros do modelo não são considerados aleatórios em si, mas passam a incorporar componentes aleatórios por meio da adição explícita de termos de efeitos aleatórios, os quais capturam a variabilidade interindividual presente nos dados.

A inclusão progressiva de termos aleatórios seguiu a estratégia descrita por Lindstrom e Bates (1990), iniciando-se com modelos contendo apenas efeitos fixos e avançando para especificações com um, dois ou três parâmetros acrescidos de efeitos aleatórios associados às unidades experimentais. A necessidade de inclusão dos termos aleatórios foi avaliada considerando a estrutura hierárquica do delineamento, a evidência de variabilidade entre réplicas e a comparação entre modelos aninhados, por meio de critérios de informação e testes da razão de verossimilhança.

Seguindo a metodologia descrita por Pinheiro e Bates (2000), as diferentes especificações dos modelos não lineares mistos foram organizadas por meio de uma notação baseada em três

índices, de forma a sistematizar simultaneamente:

- (i) combinação de parâmetros que incorporou termos de efeitos aleatórios;
- (ii) estrutura atribuída à matriz de variância-covariância dos efeitos aleatórios;
- (iii) estrutura de covariância residual adotada para as observações intraunidade, e possível heterogeneidade de variâncias entre unidades experimentais.

Assim, cada ajuste foi denotado por $\text{MOD}_{k.d.r}$, em que MOD identifica os modelos de Ørskov e McDonald, Gompertz ou Logístico; o índice k indica a estrutura hierárquica associada aos efeitos aleatórios, isto é, quais parâmetros receberam termos b_{ki} (Tabela 2); o índice d representa a forma assumida pela matriz $\mathbf{D}$ de variância-covariância dos efeitos aleatórios; e o índice r indica a estrutura especificada para a matriz residual $\mathbf{R}_i$. Dessa forma, a notação adotada permite identificar, de maneira direta, a combinação entre aleatorização dos parâmetros e as estruturas de dependência intraunidade consideradas em cada ajuste.

Em todos os modelos ajustados, os termos de efeitos aleatórios foram especificados conforme a estrutura apresentada na Tabela 2. O termo $b_{1i} \sim N(0, \sigma_{b_1}^2)$ representa a variabilidade interindividual associada ao parâmetro β_1 , relacionado ao potencial máximo de produção de gás. De forma análoga, $b_{2i} \sim N(0, \sigma_{b_2}^2)$ está associado ao parâmetro de deslocamento temporal β_2 , enquanto $b_{3i} \sim N(0, \sigma_{b_3}^2)$ representa a variabilidade associada à taxa de produção de gás β_3 .

O erro residual foi assumido como $\varepsilon_{ij} \sim N(0, \sigma_\varepsilon^2)$, mantendo independência em relação aos efeitos aleatórios. Essa especificação permite a separação explícita entre a variabilidade intra e interindividual presente nos dados longitudinais de produção de gases *in vitro*.

Tabela 2 – Modelos não lineares mistos com inclusão explícita de termos de efeitos aleatórios nos parâmetros.

Modelo	Efeitos aleatórios em:	Formulação do modelo
Modelo de Ørskov e McDonald		
OR ₁	β_1	$y_{ij} = (\beta_1 + b_{1i}) + \beta_2 [1 - \exp(-\beta_3 t_{ij})] + \varepsilon_{ij}$
OR ₂	β_2	$y_{ij} = \beta_1 + (\beta_2 + b_{2i}) [1 - \exp(-\beta_3 t_{ij})] + \varepsilon_{ij}$
OR ₃	β_3	$y_{ij} = \beta_1 + \beta_2 [1 - \exp(-(\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
OR ₄	β_1, β_2	$y_{ij} = (\beta_1 + b_{1i}) + (\beta_2 + b_{2i}) [1 - \exp(-\beta_3 t_{ij})] + \varepsilon_{ij}$
OR ₅	β_1, β_3	$y_{ij} = (\beta_1 + b_{1i}) + \beta_2 [1 - \exp(-(\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
OR ₆	β_2, β_3	$y_{ij} = \beta_1 + (\beta_2 + b_{2i}) [1 - \exp(-(\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
OR ₇	$\beta_1, \beta_2, \beta_3$	$y_{ij} = (\beta_1 + b_{1i}) + (\beta_2 + b_{2i}) [1 - \exp(-(\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
Modelo de Gompertz		
GO ₁	β_1	$y_{ij} = (\beta_1 + b_{1i}) + \beta_2 \exp[-\exp(1 - \beta_3 t_{ij})] + \varepsilon_{ij}$
GO ₂	β_2	$y_{ij} = \beta_1 + (\beta_2 + b_{2i}) \exp[-\exp(1 - \beta_3 t_{ij})] + \varepsilon_{ij}$
GO ₃	β_3	$y_{ij} = \beta_1 + \beta_2 \exp[-\exp(1 - (\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
GO ₄	β_1, β_2	$y_{ij} = (\beta_1 + b_{1i}) + (\beta_2 + b_{2i}) \exp[-\exp(1 - \beta_3 t_{ij})] + \varepsilon_{ij}$
GO ₅	β_1, β_3	$y_{ij} = (\beta_1 + b_{1i}) + \beta_2 \exp[-\exp(1 - (\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
GO ₆	β_2, β_3	$y_{ij} = \beta_1 + (\beta_2 + b_{2i}) \exp[-\exp(1 - (\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
GO ₇	$\beta_1, \beta_2, \beta_3$	$y_{ij} = (\beta_1 + b_{1i}) + (\beta_2 + b_{2i}) \exp[-\exp(1 - (\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
Modelo Logístico		
LO ₁	β_1	$y_{ij} = (\beta_1 + b_{1i}) + \beta_2 / [1 + \exp(2 - 4\beta_3 t_{ij})] + \varepsilon_{ij}$
LO ₂	β_2	$y_{ij} = \beta_1 + (\beta_2 + b_{2i}) / [1 + \exp(2 - 4\beta_3 t_{ij})] + \varepsilon_{ij}$
LO ₃	β_3	$y_{ij} = \beta_1 + \beta_2 / [1 + \exp(2 - 4(\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
LO ₄	β_1, β_2	$y_{ij} = (\beta_1 + b_{1i}) + (\beta_2 + b_{2i}) / [1 + \exp(2 - 4\beta_3 t_{ij})] + \varepsilon_{ij}$
LO ₅	β_1, β_3	$y_{ij} = (\beta_1 + b_{1i}) + \beta_2 / [1 + \exp(2 - 4(\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
LO ₆	β_2, β_3	$y_{ij} = \beta_1 + (\beta_2 + b_{2i}) / [1 + \exp(2 - 4(\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$
LO ₇	$\beta_1, \beta_2, \beta_3$	$y_{ij} = (\beta_1 + b_{1i}) + (\beta_2 + b_{2i}) / [1 + \exp(2 - 4(\beta_3 + b_{3i}) t_{ij})] + \varepsilon_{ij}$

Nota: OR = modelo de Ørskov e McDonald; GO = modelo de Gompertz; LO = modelo Logístico. Os termos $b_{ki} \sim N(0, \sigma_{b_k}^2)$ representam os efeitos aleatórios associados aos parâmetros β_k , enquanto $\varepsilon_{ij} \sim N(0, \sigma_\varepsilon^2)$ denota o erro residual.

A estimação dos parâmetros de efeitos fixos e dos componentes de variância foi realizada por meio dos métodos de máxima verossimilhança (MV) e máxima verossimilhança restrita (MVR), conforme descrito por Pinheiro e Bates (2000). Essas abordagens foram empregadas tanto no ajuste dos modelos quanto na comparação entre diferentes estruturas hierárquicas, permitindo a seleção do modelo mais parcimonioso com base em critérios de informação, como o AIC e o BIC.

Seguindo a metodologia proposta por Pinheiro e Bates (2000), foram avaliadas, para cada modelo não linear, três estruturas para a matriz de variância-covariância dos efeitos aleatórios **D**: (i) estrutura identidade (pdIdent, MOD_{k,1,r}), assumindo variância comum entre os parâmetros aleatórios; (ii) estrutura diagonal heterocedástica (pdDiag, MOD_{k,2,r}), permitindo variâncias distintas e independentes; e (iii) estrutura simétrica não estruturada (pdSymm, MOD_{k,3,r}), que admite variâncias e covariâncias livres entre os parâmetros. Essa estratégia possibilita avaliar diferentes padrões de variabilidade interindividual, evitando tanto a sub quanto a

superparametrização do modelo.

A comparação entre as estruturas de efeitos aleatórios foi conduzida por meio de testes de razão de verossimilhança (TRV) e critérios de informação, respeitando-se a hierarquia entre modelos encaixados. Nesta etapa inicial, a presença de autocorrelação residual foi investigada por meio da função de autocorrelação (ACF), com base nos resíduos condicionais, a fim de avaliar a adequação da hipótese de independência residual antes da especificação da matriz $\mathbf{R}$.

No processo subsequente, foi avaliada a estrutura de variância-covariância residual $\mathbf{R}$, responsável por representar a dependência temporal entre as medições realizadas dentro da mesma unidade experimental, bem como possíveis heterogeneidades de variância entre unidades. Após a definição da estrutura de efeitos aleatórios (índice k) e da forma da matriz $\mathbf{D}$ (índice d), foram comparadas diferentes especificações para a matriz residual $\mathbf{R}$ (índice r).

Consideraram-se as estruturas de correlação de simetria composta (CS), exponencial (EXP) e gaussiana (GAUS), cada uma avaliada sob duas especificações: com variância residual homogênea e com heterogeneidade de variâncias modelada por meio da função `varIdent()` (ID). Dessa forma, foram ajustados seis modelos candidatos, denotados por $\text{MOD}_{k,d,1}$ a $\text{MOD}_{k,d,6}$, correspondendo às combinações entre as três estruturas de correlação e a presença ou ausência da função de variância `varIdent()`.

A seleção da estrutura residual foi conduzida de forma hierárquica, mantendo-se fixas as especificações previamente definidas para $\mathbf{D}$ e para os termos de efeitos aleatórios, e comparando-se as diferentes estruturas candidatas para $\mathbf{R}$. A comparação entre modelos foi realizada com base nos critérios de informação de Akaike (AIC) e Bayesiano (BIC), bem como no valor do logaritmo da verossimilhança ($\log\text{Lik}$), e sua adequação por meio da análise dos resíduos condicionais padronizados, utilizando envelopes simulados baseados em gráficos half-normal (HNP). Essa abordagem permite avaliar a compatibilidade dos resíduos com a distribuição assumida pelo modelo, sendo particularmente informativa em contextos de modelos não lineares e confirmando a escolha estrutural dos modelos adotada.

A capacidade preditiva dos MNLM ajustados, com especificação dos parâmetros incluindo efeitos aleatórios e estruturas de variância-covariância para os efeitos aleatórios e residuais, foi avaliada com base em um conjunto de dados de teste, distinto daquele utilizado no processo de estimação dos parâmetros. Sendo conduzida em dois níveis de predição: marginal (nível 0) e condicional (nível 1). No nível marginal, as predições foram obtidas exclusivamente a partir dos efeitos fixos do modelo, refletindo o comportamento médio da população. No nível condicional, incorporaram-se os efeitos aleatórios estimados, permitindo capturar a variabilidade específica de cada unidade experimental e, conseqüentemente, obter predições mais ajustadas às trajetórias individuais.

As métricas preditivas utilizadas incluíram a raiz do erro quadrático médio (RMSE), o erro absoluto médio (MAE), o erro percentual absoluto médio (MAPE), o coeficiente de

correlação de Pearson (Corr), o coeficiente de concordância de Lin (CCC), apresentados na Tabela 1, e gráficos de valores preditos versus valores observados, permitindo avaliar diferentes aspectos do desempenho preditivo, incluindo a magnitude dos erros, a proporcionalidade das discrepância e o grau de associação e concordância entre valores observados e preditos sob os dois níveis de predição, contribuindo para a seleção de modelos mais robustos e biologicamente coerentes na descrição da cinética de produção de gases.

Com a definição do modelo final, procedeu-se à comparação entre tratamentos por meio do teste de Tukey aplicado às médias marginais estimadas a partir do modelo misto não linear selecionado, garantindo que as comparações respeitassem a estrutura hierárquica e a dependência longitudinal dos dados. Sendo possível a partir dos parâmetros estimados do modelo calcular a degradabilidade efetiva (DE) segundo a formulação proposta por Ørskov e McDonald (1979) (Equação 1.4), considerando taxa de passagem ruminal fixa de $k = 0,05 \text{ h}^{-1}$. A DE foi utilizada como métrica integrada do potencial fermentativo dos substratos, permitindo a hierarquização das amostras sob uma perspectiva biológica e quantitativa.

Com base na síntese dos resultados do ajuste estatístico, das análises diagnósticas, das métricas de desempenho preditivo e da interpretação biológica dos parâmetros, foi identificado o modelo com melhor desempenho global e os tratamentos com maior potencial fermentativo.

RESULTADOS E DISCUSSÕES

3.1 ANÁLISE EXPLORATÓRIA

Os dados analisados neste estudo correspondem à produção acumulada de gás *in vitro*, registrada ao longo do tempo a partir da incubação de amostras de coprodutos sob condições controladas. As medições foram realizadas em intervalos regulares de 12 minutos, durante um período total de 72 horas, resultando em um número elevado de observações por curva.

Cada unidade experimental foi definida pela combinação entre réplica e amostra, sendo associada a uma curva de produção de gás ao longo do tempo. Dessa forma, as observações possuem estrutura hierárquica, na qual medidas repetidas no tempo estão aninhadas dentro de cada unidade experimental. No contexto dos modelos não lineares mistos, essa estrutura é incorporada por meio da inclusão de efeitos aleatórios associados às unidades experimentais, permitindo capturar a variabilidade entre curvas, enquanto a dependência entre observações ao longo do tempo dentro de cada unidade é modelada por meio de estruturas adequadas para a matriz de covariância residual.

A inspeção gráfica preliminar dos dados (Figura 2) evidenciou um padrão geral compatível com o comportamento esperado para a produção acumulada de gás. No entanto, também foram identificadas curvas com trechos irregulares, interrupções abruptas, baixa produção inicial de gás ou oscilações incompatíveis com a dinâmica fermentativa típica. Exemplos representativos desse comportamento podem ser observados na réplica 2 da amostra 9 e nas réplicas 1 e 3 da amostra 14.

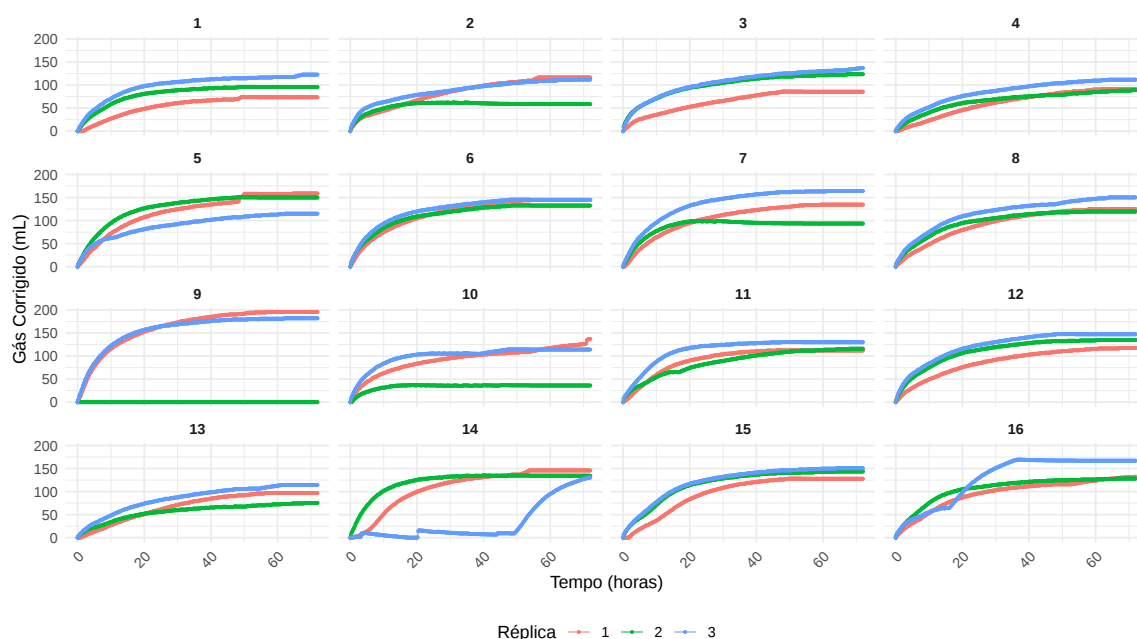


Figura 2 – Produção de gás corrigido ao longo do tempo para as amostras.

Diante dessas inconsistências, adotou-se o procedimento exploratório de seleção e limpeza dos dados, baseado na inspeção detalhada das curvas individuais (Apendice A). Optou-se pela exclusão integral da curva quando a confiabilidade das observações se encontrava substancialmente comprometida.

A avaliação das curvas individuais revelou unidades experimentais com comportamento incompatível com o padrão monotonicamente crescente esperado para medidas acumuladas de fermentação ruminal. Entre os principais problemas observados destacam-se: produção de gás praticamente constante desde tempos iniciais de incubação, valores finais substancialmente inferiores ao conjunto majoritário das amostras e oscilações abruptas ao longo do tempo. Tais padrões sugerem possíveis falhas técnicas durante o processo de incubação ou de leitura dos volumes produzidos.

Conforme discutido anteriormente na Seção 1.1.1, o ajuste de modelos não lineares a dados de produção acumulada é sensível à presença de observações inconsistentes, uma vez que esses modelos assumem trajetórias suaves e crescentes ao longo do tempo. Desvios desse comportamento podem comprometer a convergência dos algoritmos de estimação e resultar em parâmetros instáveis ou biologicamente implausíveis, especialmente em modelos não lineares mistos com estrutura hierárquica (PINHEIRO; BATES, 2000).

A exclusão dessas unidades experimentais teve como objetivo preservar a estabilidade dos ajustes, assegurar a plausibilidade biológica dos parâmetros estimados e garantir que as inferências refletissem adequadamente a dinâmica fermentativa dos coprodutos avaliados. No total, foram excluídas cinco curvas individuais de produção acumulada de gás, correspondentes às

combinações específicas de réplica e tratamento: (1:3), (2:9), (2:10), (3:14) e (3:16) apresentados na Figura 3.

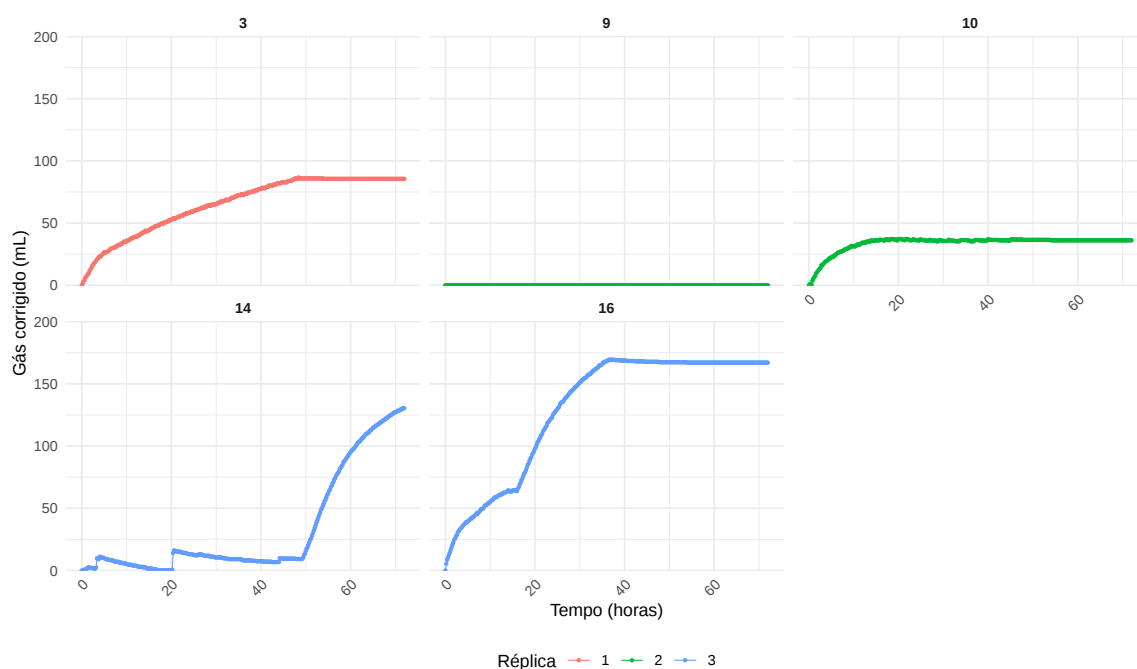


Figura 3 – Curvas de produção de gás excluídas das análises.

Para as demais unidades experimentais, adotou-se um procedimento de truncamento temporal das curvas, no qual o intervalo de observação foi individualmente restringido para cada combinação de réplica e tratamento. Os limites superiores de tempo foram definidos com base em inspeção gráfica detalhada das curvas de produção acumulada de gás, variando entre 27 e 72 horas de incubação. Nos casos em que não se observaram desvios relevantes, considerou-se a totalidade do período experimental.

A aplicação desses cortes justificou-se pela identificação de padrões atípicos na fase final de incubação, tais como perda de monotonicidade, oscilações não compatíveis com o processo biológico ou aumento sistemático dos resíduos, indicando deterioração da qualidade de ajuste do modelo não linear misto (Apêndice A). Nessas situações, a inclusão dessas observações poderia introduzir viés na estimação dos parâmetros e comprometer a interpretação biológica dos resultados.

Dessa forma, o intervalo de observação de cada curva foi restrito aos trechos considerados biologicamente plausíveis e estatisticamente informativos, contribuindo para maior robustez dos ajustes e consistência das inferências. Ressalta-se que a definição dos pontos de truncamento foi realizada por meio de inspeção minuciosa, curva a curva, sendo os respectivos gráficos apresentados no Apêndice A, assegurando transparência e reprodutibilidade ao processo de seleção dos dados.

Após a etapa de inspeção exploratória e exclusão das curvas consideradas inconsistentes,

procedeu-se à reorganização do conjunto de dados com base em um conjunto discreto de tempos de incubação, selecionados conforme o protocolo experimental descrito por Cabral et al. (2019). Esses tempos correspondem a instantes específicos de leitura ao longo da cinética de produção de gás, definidos como: 0, 2, 4, 6, 8, 10, 12, 24, 26, 28, 30, 32, 36, 48, 52, 54, 56, 58, 60, 64 e 72 horas.

A adoção desses tempos fundamenta-se tanto em critérios experimentais quanto estatísticos. Do ponto de vista biológico, buscou-se contemplar adequadamente as diferentes fases do processo fermentativo, incluindo o período inicial de colonização microbiana, a fase de maior atividade fermentativa e a aproximação assintótica da produção acumulada de gás. Conforme discutido por Teixeira et al. (2016), esquemas temporais com maior concentração de leituras nas horas iniciais e observações adicionais em tempos tardios tendem a favorecer a identificação dos parâmetros cinéticos e a adequada descrição da curva de fermentação.

Entretanto, a utilização de conjuntos temporais muito reduzidos, com número insuficiente de observações ao longo da curva, pode comprometer a identificabilidade dos parâmetros não lineares, sobretudo em modelos com múltiplos componentes cinéticos. A escassez de pontos observacionais dificulta a caracterização simultânea da fase inicial, da região de maior crescimento e do comportamento assintótico, podendo resultar em instabilidade numérica e dificuldades de convergência dos algoritmos iterativos de estimação. Dessa forma, optou-se pela manutenção dos tempos propostos por Cabral et al. (2019), os quais oferecem quantidade de observações suficiente e distribuição temporal adequada para representar a dinâmica fermentativa sem perda substancial de informação.

Além disso, a manutenção dos tempos propostos por Cabral et al. (2019) permitiu preservar a coerência com o delineamento experimental originalmente adotado, assegurando comparabilidade metodológica entre os resultados obtidos e aqueles reportados no estudo de referência. Do ponto de vista estatístico, a utilização de tempos de observação previamente definidos contribui para a estabilidade numérica dos algoritmos de estimação em modelos não lineares de efeitos mistos, além de favorecer a comparabilidade entre réplicas e tratamentos, uma vez que garante uma estrutura temporal comum entre as unidades experimentais, aspecto importante para a adequada especificação e estimação desses modelos (PINHEIRO; BATES, 2000).

A seleção dos tempos foi implementada por meio da filtragem direta do banco de dados, mantendo-se apenas as observações correspondentes aos instantes de incubação previamente especificados. Como consequência das etapas de depuração das curvas e da restrição aos tempos selecionados, o conjunto de dados final apresentou redução em relação à base original. Inicialmente, o banco continha 17328 observações, sendo reduzido para 715 observações após os procedimentos de exclusão e filtragem temporal.

Na Figura 4 são apresentadas as curvas de produção de gás por réplica após o processo de depuração e seleção dos tempos de incubação. Essa visualização permite avaliar a consis-

tência intra-amostral das trajetórias fermentativas e fornece uma base adequada para os ajustes subsequentes dos modelos não lineares de efeitos mistos.

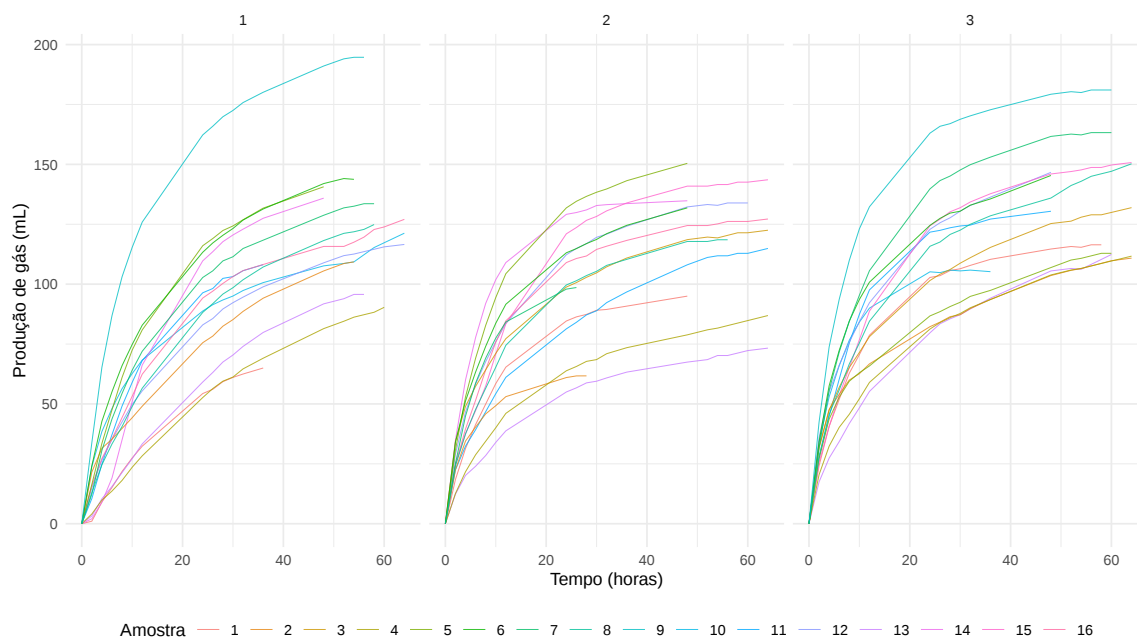


Figura 4 – Curvas de produção de gás por réplica que permaneceram para ajuste (após exclusão de amostras inconsistentes)

Após a remoção das unidades experimentais classificadas como inconsistentes, procedeu-se à restrição do conjunto de dados aos tempos de incubação selecionados para a modelagem, a Tabela 3 apresenta as estatísticas descritivas da produção de gás antes e após a “limpeza” de dados e filtragem temporal.

Tabela 3 – Estatísticas descritivas gerais da produção de gás (G) nos bancos completo e filtrado

Estatística	Completo (mL)	Filtrado (mL)
1º Quartil	66,76	53,02
Mediana	102,44	89,44
Média	95,60	84,92
3º Quartil	126,67	118,01
Desvio padrão	43,03	44,31

A restrição do conjunto de dados, decorrente da exclusão de curvas inconsistentes e da seleção de tempos de incubação específicos, implicou alterações na distribuição da produção de gás observada. De modo geral, observa-se modificação nas medidas de posição (quartis, mediana e média), refletindo a redefinição do domínio amostral considerado na análise.

Essas mudanças estão diretamente associadas à remoção de observações provenientes de trechos finais das curvas ou de unidades experimentais com comportamento atípico, os quais,

em geral, contribuem para maior variabilidade e potencial distorção na distribuição da variável resposta.

No que se refere à dispersão, o desvio padrão apresenta variação após o processo de filtragem, indicando alteração na heterogeneidade dos dados. Esse resultado é consistente com a exclusão de porções das curvas com menor aderência ao modelo, o que tende a produzir um conjunto de dados mais homogêneo e estatisticamente coerente com as suposições do modelo não linear misto.

Ressalta-se que tais modificações não configuram perda de informação estatística relevante, mas sim uma redefinição controlada do suporte temporal e experimental da análise. Esse procedimento é fundamental para assegurar maior estabilidade numérica nos algoritmos de estimação, bem como melhorar a identificabilidade dos parâmetros estruturais do modelo, conforme discutido em Pinheiro e Bates (2000).

3.2 AJUSTE DE MODELOS NÃO LINEARES

Com base nas observações experimentais da produção acumulada de gás ao longo do tempo (Figura 4), procedeu-se inicialmente ao ajuste de três modelos não lineares amplamente empregados para descrever a cinética de fermentação ruminal: o modelo de Ørskov e McDonald (1979) (OR), o modelo de Gompertz (GO) e o modelo Logístico (LO), conforme discutido na Seção 2.2. Para padronização da escrita, os parâmetros desses modelos foram denotados por β_1 , β_2 e β_3 .

De modo geral, os parâmetros podem ser interpretados como componentes complementares da dinâmica de produção de gás: β_1 está associado ao nível inicial da resposta (produção instantânea no tempo $t = 0$, quando aplicável); β_2 representa a fração potencialmente fermentável, de forma que o potencial assintótico de produção de gás é dado por $\beta_1 + \beta_2$ e β_3 corresponde ao parâmetro de velocidade, relacionado à taxa com que a produção acumulada se aproxima do platô.

Foram adotados valores iniciais plausíveis para cada parâmetro, com base no conhecimento biológico prévio acerca da cinética fermentativa e nas recomendações da literatura especializada. Após o ajuste dos modelos na forma fixa, foram extraídas as estimativas individuais dos parâmetros β_1 , β_2 e β_3 , calculando-se, para cada modelo, a média, o desvio-padrão (DP) e o coeficiente de variação (CV), como medidas descritivas da variabilidade entre curvas.

A Tabela 4 apresenta o resumo dessas estimativas considerando as unidades experimentais efetivamente utilizadas após a exclusão de curvas inconsistentes e aplicação dos cortes temporais. Embora o delineamento experimental contemplasse inicialmente até 48 curvas por modelo (16 amostras em três réplicas), os resultados apresentados refletem apenas as curvas válidas para ajuste.

Observa-se que o parâmetro β_2 , associado ao potencial de produção de gás (ou nível assintótico da curva, conforme a parametrização adotada), apresentou médias relativamente próximas entre os três modelos, variando de 111,83 a 121,20. Os coeficientes de variação situaram-se entre 21,75% e 23,42%, indicando comportamento estável e baixa sensibilidade desse parâmetro à escolha da função não linear. Esse resultado sugere consistência na estimativa do potencial máximo de produção entre as diferentes especificações modeladas.

O parâmetro β_3 , relacionado à taxa de produção de gás, apresentou coeficientes de variação entre 38,08% e 45,18%, evidenciando variabilidade moderada entre as curvas individuais. O modelo de Ørskov (OR) apresentou maior dispersão relativa, enquanto os modelos Gompertz (GO) e Logístico (LO) apresentaram maior estabilidade nesse parâmetro. Esse comportamento é compatível com a natureza dos experimentos de fermentação *in vitro*, nos quais diferenças na degradabilidade do substrato e na atividade microbiana influenciam diretamente a velocidade do processo fermentativo.

Em contraste, o parâmetro β_1 apresentou elevada variabilidade relativa, com coeficientes de variação de 130,16% no modelo Ørskov (OR), 118,15% no modelo Gompertz (GO) e 393,23% no modelo Logístico (LO). Destaca-se, em particular, o comportamento do modelo Logístico, que apresentou variabilidade extremamente elevada nesse parâmetro, indicando possível instabilidade na sua estimação ou elevada sensibilidade à variabilidade entre unidades experimentais. Esse resultado reflete tanto diferenças estruturais entre curvas individuais quanto limitações associadas à parametrização do modelo para descrever a fase inicial da produção de gás.

Ainda que elevados, tais coeficientes de variação não indicam necessariamente falha de estimação, mas sim maior instabilidade relativa inerente à forma funcional do parâmetro em cada modelo, especialmente no caso do modelo Logístico.

Tabela 4 – Comparação descritiva dos parâmetros estimados nos modelos não lineares fixos

Modelo	Estatística	β_1	β_2	β_3
OR	Média	4,277	121,20	0,0765
	DP	5,566	26,36	0,0346
	CV (%)	130,16	21,75	45,18
GO	Média	5,967	111,83	0,1510
	DP	7,050	25,37	0,0614
	CV (%)	118,15	22,69	40,67
LO	Média	1,650	113,27	0,0576
	DP	6,489	26,53	0,0219
	CV (%)	393,23	23,42	38,08

Nota: OR – modelo de Ørskov & McDonald (1979); GO – modelo de Gompertz; LO – modelo Logístico. DP: desvio-padrão; CV: coeficiente de variação.

Esses resultados reforçam a necessidade de uma avaliação inferencial mais detalhada, baseada em intervalos de confiança, a fim de complementar a análise descritiva e permitir uma interpretação mais robusta da precisão e da incerteza associadas às estimativas dos parâmetros nos diferentes modelos.

Nos intervalos de confiança dos parâmetros estimados para cada amostra nos modelos de Ørskov, Gompertz e Logístico (Figuras 5, 6 e 7), os parâmetros são interpretados de acordo com a parametrização adotada nos modelos não lineares. Observa-se que, independentemente do modelo considerado, há ampla dispersão nas estimativas entre as diferentes curvas, refletindo a heterogeneidade na dinâmica de produção de gás.

Nota-se, que muitos dos intervalos não se sobrepõem, essa ausência de sobreposição evidencia diferenças estatisticamente significativas entre as amostras, o que justifica a modelagem dos parâmetros como efeitos aleatórios nos modelos mistos.

Dessa forma, os gráficos reforçam a necessidade do uso de modelos não lineares mistos, uma vez que os modelos fixos não capturam adequadamente a variabilidade individual presente nos dados, negligenciando a estrutura hierárquica e a correlação intraunidade típica de experimentos longitudinais. A inclusão de efeitos aleatórios permite modelar diretamente a variação entre unidades experimentais, fornecendo estimativas mais realistas dos parâmetros populacionais e melhorando a qualidade do ajuste (PINHEIRO; BATES, 2000).

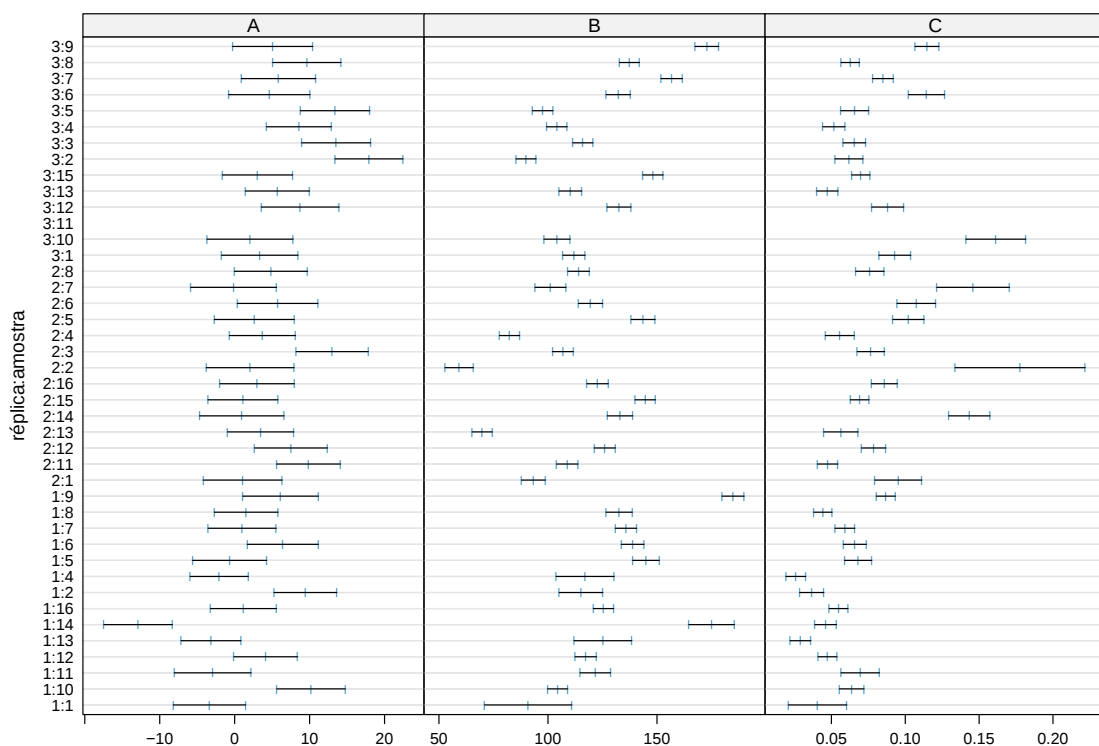


Figura 5 – Intervalos de confiança (95%) dos parâmetros estimados para cada amostra, modelo de Ørskov & McDonald (OR). ($A = \beta_1$, $B = \beta_2$, $C = \beta_3$).

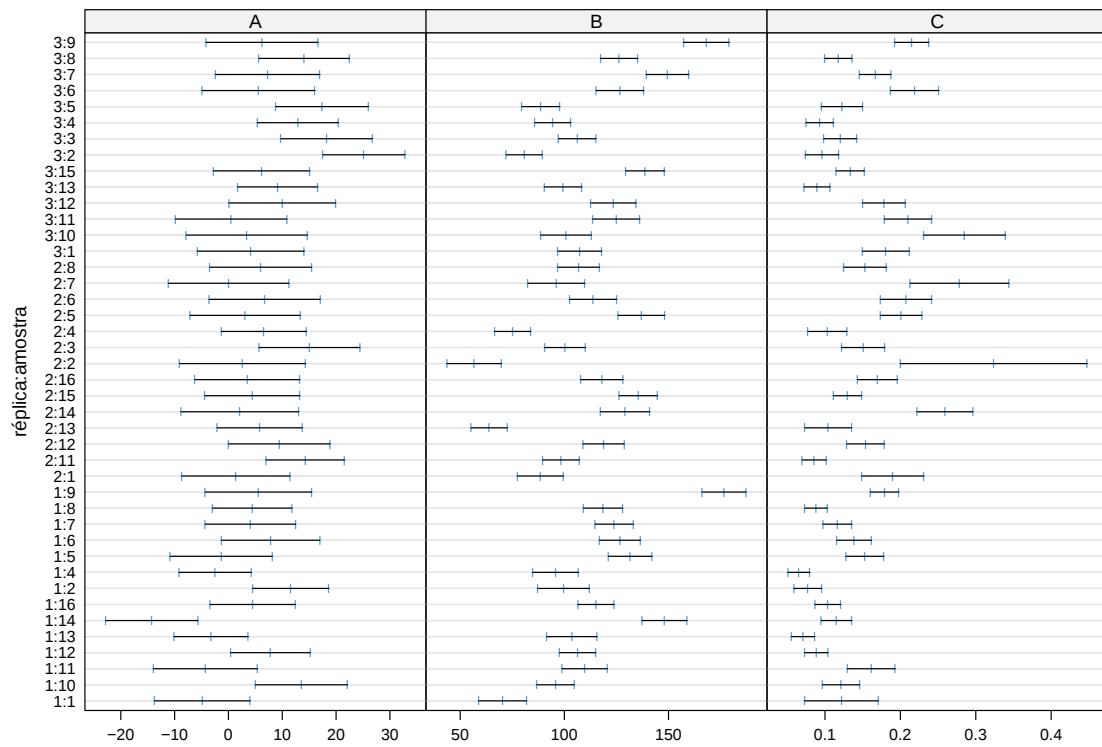


Figura 6 – Intervalos de confiança (95%) dos parâmetros estimados para cada amostra, modelo de Gompertz (GO). ($A = \beta_1$, $B = \beta_2$, $C = \beta_3$).

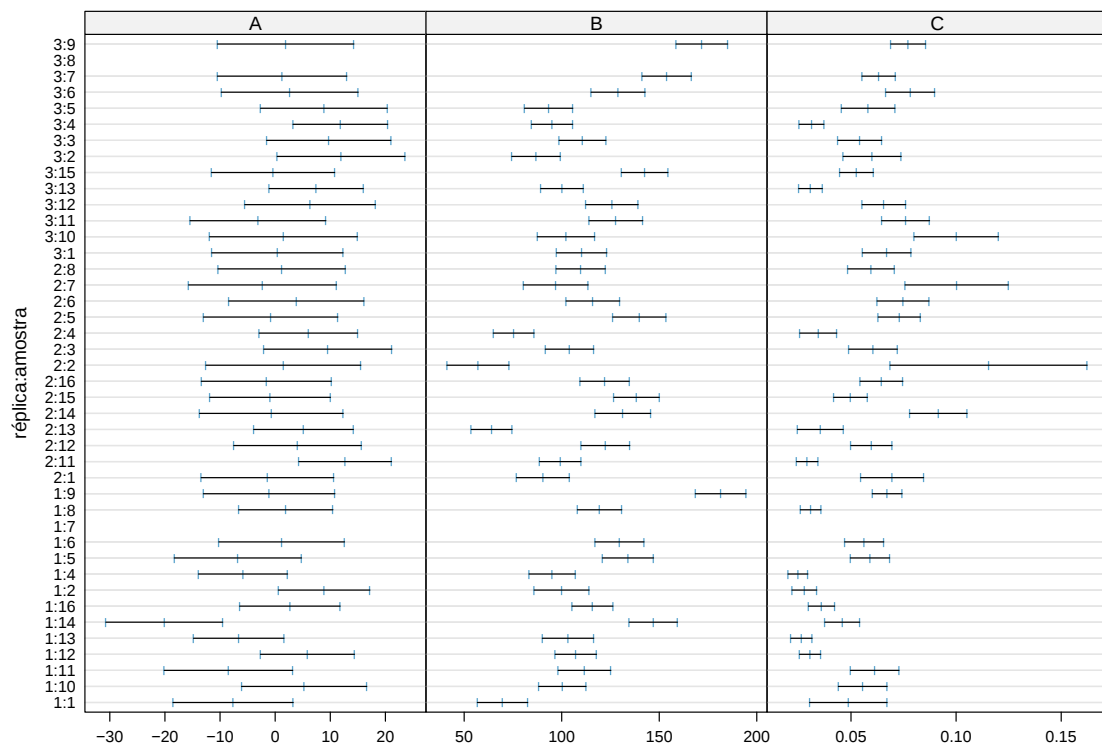


Figura 7 – Intervalos de confiança (95%) dos parâmetros estimados para cada amostra, modelo Logístico (LO). ($A = \beta_1$, $B = \beta_2$, $C = \beta_3$).

Os gráficos de intervalos de confiança obtidos a partir dos ajustes realizados por meio da função `nlsList()`, implementada no pacote `nlme` (PINHEIRO; BATES; R Core Team, 2025), (Figuras 5, 6 e 7) indicam que, para algumas curvas individuais, os intervalos associados a determinados parâmetros se estendem até valores negativos. Esse comportamento é compatível com as limitações conhecidas dos intervalos de confiança baseados em aproximações normais em modelos não lineares, sobretudo em situações nas quais os parâmetros apresentam assimetria, forte dependência entre si ou fraca identificabilidade ao nível individual (WEST; WELCH; GALECKI, 2014).

Além disso, o ajuste não linear realizado curva a curva por mínimos quadrados está sujeito a instabilidades numéricas quando a informação contida na curva é limitada, quando há colinearidade local entre parâmetros ou dificuldades de convergência. Nessas situações, as aproximações utilizadas para a obtenção de erros-padrão e intervalos de confiança podem se tornar pouco informativas, sendo comum a ocorrência de estimativas biologicamente pouco plausíveis ou mesmo falhas de convergência em algumas unidades experimentais. Conforme discutido por Pinheiro e Bates (2000), o próprio `nlsList()` pode prosseguir com o ajuste e retornar valores NA para coeficientes associados a curvas que não convergiram adequadamente.

Esses resultados devem ser interpretados no contexto exploratório do ajuste por curvas individuais, cujo objetivo principal é avaliar a variabilidade inicial dos parâmetros e diagnosticar potenciais dificuldades de identificação. Nesse sentido, os coeficientes de variação apresentados na Tabela 4 quantificam a variabilidade relativa das estimativas entre curvas, enquanto os intervalos de confiança ilustram a incerteza associada à estimação de cada curva isoladamente. Essas medidas capturam aspectos distintos da variabilidade paramétrica e, portanto, não devem ser interpretadas de forma equivalente. Em particular, parâmetros com média populacional reduzida podem apresentar coeficientes de variação elevados, mesmo quando os intervalos de confiança individuais são relativamente estreitos, reforçando a necessidade de modelagem conjunta por meio de modelos não lineares mistos.

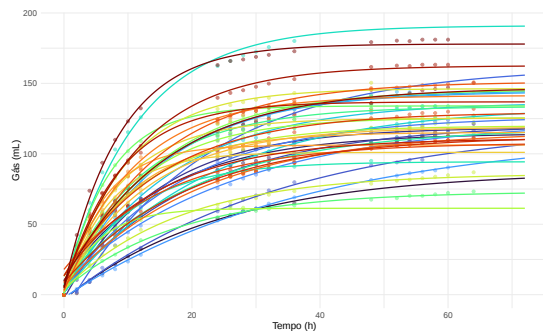
A etapa seguinte consistiu em avançar para o ajuste de modelos não lineares mistos via `nlme` (PINHEIRO; BATES; R Core Team, 2025), seguindo a formulação padrão em que os efeitos fixos β representam os valores médios populacionais dos parâmetros, enquanto os efeitos aleatórios $\mathbf{b}_i$ representam os desvios de cada unidade experimental em relação a essa média. Assume-se independência dos vetores de efeitos aleatórios entre unidades e independência dos erros intraunidade, bem como independência entre erros e efeitos aleatórios. (PINHEIRO; BATES, 2000)

Como no presente estudo o número de unidades experimentais válidas é grande em relação ao número de efeitos aleatórios do modelo, adotou-se, como ponto de partida, uma matriz $\mathbf{D}$ geral positiva-definida para a estrutura de variância-covariância dos efeitos aleatórios, conforme recomendado por Pinheiro e Bates (2000) em situações análogas. Essa escolha permite que os dados indiquem (i) a magnitude da variabilidade entre unidades e (ii) se há correlações

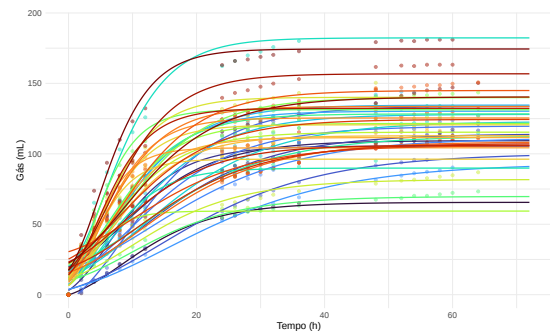
relevantes entre os componentes aleatórios; posteriormente, caso as correlações estimadas se mostrem pequenas, estruturas mais parcimoniosas (e.g., diagonal) podem ser consideradas sem perda substancial de ajuste.

A Figura 8 apresenta os ajustes individuais dos modelos não lineares de Ørskov & McDonald, Gompertz e Logístico aplicados a cada amostra experimental (réplica:tratamento). As curvas foram obtidas a partir do ajuste por amostra com o método `nlsList()`, permitindo visualizar a variabilidade entre os perfis de fermentação.

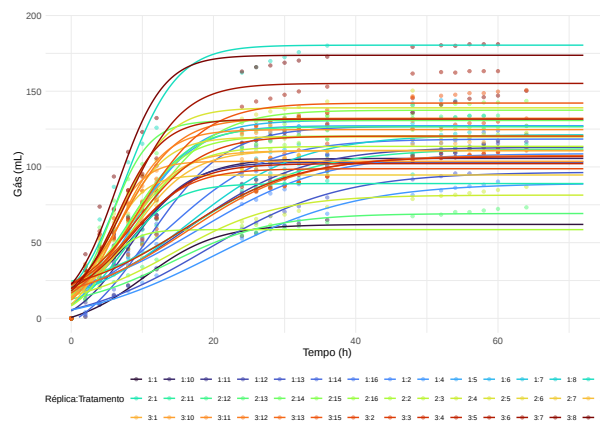
Observa-se que as curvas seguem padrões distintos de produção de gás, algumas apresentam crescimento mais rápido, enquanto outras mostram uma taxa mais lenta de fermentação ou um volume final de gás menor. Essa variabilidade reforça que os parâmetros do modelo se comportam de maneira diferente entre as unidades experimentais evidenciando a limitação dos modelos fixos em representar adequadamente a estrutura hierárquica dos dados (Figura 8). A adoção de modelos não lineares mistos torna-se, portanto, essencial para capturar a heterogeneidade entre amostras. Essa abordagem permite incorporar efeitos aleatórios nos parâmetros dos modelos, melhorando a capacidade preditiva e a precisão das inferências estatísticas (PINHEIRO; BATES, 2000).



(a) Modelo de Ørskov & McDonald (1979)



(b) Modelo de Gompertz



(c) Modelo Logístico

Figura 8 – Curvas ajustadas para cada amostra segundo os modelos não lineares fixos de Ørskov, Gompertz e Logístico.

Diferentemente do modelo de Ørskov & McDonald, os modelos de Gompertz e Logístico apresentam valores iniciais de produção de gás próximos de zero. Esse comportamento decorre da própria forma funcional sigmoideal dessas formulações, nas quais a produção no instante inicial não é controlada diretamente por um parâmetro associado à fração imediatamente solúvel, mas resulta da combinação entre os parâmetros estruturais da curva e da forma matemática da função.

Do ponto de vista experimental, observa-se que as curvas individuais de produção acumulada de gás não iniciam necessariamente em zero, havendo uma faixa relativamente ampla de valores iniciais. Essa variabilidade pode estar associada a fatores experimentais, como pequenas diferenças no momento da leitura inicial, ajustes de calibração do sistema de mensuração ou mesmo à presença de frações prontamente fermentáveis que iniciam a produção de gás antes do primeiro registro temporal e corrobora com as informações apresentadas na Tabela 4.

No modelo de Ørskov & McDonald, essa produção inicial é capturada explicitamente pelo parâmetro β_1 , o qual representa a fração imediatamente solúvel do substrato. Já nos modelos de Gompertz e Logístico, a condição inicial próxima de zero é uma implicação direta da parametrização adotada. Assim, a ausência de intercepto positivo nesses modelos não deve ser interpretada como inadequação do ajuste, mas como característica estrutural da função matemática empregada.

3.3 AJUSTE DOS MODELOS NÃO LINEARES MISTOS

Dando continuidade à análise, foram ajustados modelos não lineares mistos com o objetivo de representar adequadamente a estrutura hierárquica dos dados e incorporar a variabilidade observada entre as unidades experimentais. A etapa anterior, baseada em modelos não lineares fixos, evidenciou que, embora as curvas médias descrevam satisfatoriamente a tendência geral da produção de gás ao longo do tempo, existe considerável heterogeneidade entre as unidades experimentais, tanto na forma quanto na taxa de crescimento das curvas.

Para acomodar essa variabilidade, adotou-se uma estrutura de efeitos mistos, permitindo a inclusão de termos aleatórios nos parâmetros dos modelos. Os efeitos fixos foram especificados em função dos tratamentos para os parâmetros β_2 e β_3 , de modo a capturar diferenças sistemáticas entre os coprodutos avaliados, enquanto o parâmetro β_1 foi considerado comum entre os tratamentos. Essa especificação foi adotada considerando aspectos biológicos e estatísticos do processo avaliado. O parâmetro β_1 , associado à fração imediatamente disponível ou à produção inicial de gás, tende a apresentar baixa variabilidade entre os tratamentos, uma vez que todas as amostras são submetidas às mesmas condições iniciais de incubação. Dessa forma, diferenças sistemáticas entre os coprodutos manifestam-se predominantemente nos parâmetros relacionados à fração potencialmente degradável e à taxa de fermentação (β_2 e β_3).

Além disso, a inclusão de efeitos fixos por tratamento apenas nesses parâmetros contribui

para uma parametrização mais parcimoniosa, evitando a introdução de efeitos desnecessários em componentes cuja variabilidade não é estruturalmente relevante. Sob o ponto de vista da modelagem, essa escolha favorece a estabilidade numérica do processo de estimação e a interpretabilidade dos parâmetros, mantendo o modelo coerente com a dinâmica biológica da produção de gás.

Os efeitos aleatórios foram associados à combinação entre réplica e tratamento. Essa escolha reflete o delineamento experimental, no qual cada curva corresponde a uma unidade experimental distinta, sujeita a variações específicas decorrentes tanto da réplica quanto da própria amostra. Dessa forma, as réplicas associadas a cada tratamento foram consideradas como o nível hierárquico no qual se introduzem os efeitos aleatórios, permitindo que os parâmetros variassem entre as curvas individuais em torno de médias populacionais definidas pelos efeitos fixos. Essa formulação está em consonância com a abordagem proposta por Pinheiro e Bates (2000), na qual os efeitos fixos representam os valores médios dos parâmetros na população, enquanto os efeitos aleatórios descrevem as devições específicas de cada unidade experimental em relação a essas médias.

Na Tabela 5 são apresentados os resultados dos ajustes dos modelos não lineares mistos de Ørskov & McDonald, Gompertz e Logístico, considerando diferentes combinações de inclusão de efeitos aleatórios nos parâmetros. Os modelos foram comparados com base nos critérios de informação de Akaike (AIC), Bayesiano (BIC) e no valor da log-verossimilhança, com o objetivo de identificar a estrutura mais parcimoniosa e adequada à variabilidade observada nos dados.

Sob a perspectiva dos critérios de informação, observa-se que, para os três modelos avaliados, a especificação que incorporou efeitos aleatórios simultaneamente nos parâmetros β_1 , β_2 e β_3 apresentou os menores valores de AIC e BIC. Esse resultado indica melhor compromisso entre qualidade de ajuste e complexidade do modelo, evidenciando que a inclusão da variabilidade entre unidades experimentais nos três componentes estruturais da curva contribui de maneira relevante para a descrição dos dados.

Comparativamente, embora a especificação com efeitos aleatórios apenas em β_2 e β_3 também tenha apresentado desempenho substancialmente superior às versões mais restritivas, a inclusão adicional do parâmetro β_1 promoveu redução adicional dos critérios de informação nos três modelos (Ørskov, Gompertz e Logístico). Tal comportamento indica que esse parâmetro também captura parte da heterogeneidade entre as curvas individuais, cuja modelagem como efeito aleatório melhora o ajuste global.

Esse achado é consistente com a análise exploratória dos intervalos de confiança obtidos a partir dos ajustes via `nlsList()` (Figuras 5, 6 e 7), os quais evidenciaram dispersão entre curvas também para o parâmetro associado ao nível inicial da resposta, notadamente nas amostras (1:1), (1:14), (1:15) e (1:10), cujas trajetórias iniciais não se assemelham às demais réplicas do experimento. Observou-se que essas curvas apresentaram acúmulo de gás substancialmente

inferior nas primeiras horas de incubação, quando comparadas tanto às demais réplicas do mesmo tratamento quanto ao padrão geral observado no conjunto de dados.

Sob o ponto de vista biológico, o parâmetro β_1 , interpretado como a fração imediatamente disponível ou produção inicial de gás, tende a apresentar variabilidade limitada entre unidades experimentais, especialmente considerando que todas as amostras iniciam o ensaio sob condições padronizadas e com volume inicial igual a zero. Dessa forma, a elevada variabilidade observada em algumas curvas específicas pode estar mais associada a particularidades experimentais ou comportamentos individuais atípicos do que a uma heterogeneidade estrutural sistemática entre tratamentos.

Tabela 5 – Critérios de seleção dos modelos não lineares mistos com diferentes especificações de efeitos aleatórios.

Modelo	k	Inclusão de efeitos aleatórios	AIC	BIC	logLik
OR _k	1	β_1	4835	4995	-2383
	2	β_2	4844	5004	-2387
	3	β_3	4723	4883	-2326
	4	$\beta_1 + \beta_2$	4622	4791	-2274
	5	$\beta_1 + \beta_3$	4649	4818	-2288
	6	$\beta_2 + \beta_3$	4332	4501	-2129
	7	$\beta_1 + \beta_2 + \beta_3$	4249	4432	-2085
GO _k	1	β_1	5282	5442	-2606
	2	β_2	5293	5453	-2612
	3	β_3	5315	5475	-2622
	4	$\beta_1 + \beta_2$	5218	5387	-2572
	5	$\beta_1 + \beta_3$	5235	5404	-2580
	6	$\beta_2 + \beta_3$	5132	5302	-2529
	7	$\beta_1 + \beta_2 + \beta_3$	5125	5308	-2523
LO _k	1	β_1	5445	5605	-2688
	2	β_2	5447	5607	-2688
	3	β_3	5512	5672	-2721
	4	$\beta_1 + \beta_2$	5408	5577	-2667
	5	$\beta_1 + \beta_3$	5436	5606	-2681
	6	$\beta_2 + \beta_3$	5361	5531	-2644
	7	$\beta_1 + \beta_2 + \beta_3$	5358	5541	-2639

Nota: AIC = Critério de Informação de Akaike; BIC = Critério de Informação Bayesiano; logLik = log-verossimilhança. OR – modelo de Ørskov e McDonald; GO – modelo de Gompertz; LO – modelo Logístico.

Estudos semelhantes conduzidos por Medeiros et al. (2020) relatam que a fração inicial tende a apresentar menor variabilidade biológica quando comparada aos parâmetros associados à fração potencialmente degradável e à taxa de degradação, reforçando a interpretação de que a aleatorização de β_1 nem sempre é estritamente necessária do ponto de vista biológico. Ainda assim, os critérios de informação indicaram que sua inclusão como efeito aleatório contribuiu

para capturar a heterogeneidade observada nas curvas individuais, especialmente diante da presença de trajetórias com acúmulo inicial reduzido.

Esse contraste entre interpretação biológica e evidência estatística reforça a importância de considerar simultaneamente critérios quantitativos de ajuste e análise exploratória das curvas individuais na especificação final do modelo misto. Conforme discutido por Pinheiro e Bates (2000), a inspeção dos intervalos individuais constitui ferramenta importante para orientar a especificação da estrutura de efeitos aleatórios em modelos não lineares mistos, especialmente quando se busca equilíbrio entre parcimônia e capacidade explicativa. De forma consistente entre as três funções avaliadas, a especificação completa com efeitos aleatórios em β_1 , β_2 e β_3 (modelo $k = 7$) apresentou o melhor desempenho segundo os critérios de informação, sendo considerada como referência na avaliação das estruturas de efeitos aleatórios.

A análise conjunta dos critérios de informação (AIC e BIC), da log-verossimilhança (logLik) e dos testes da razão de verossimilhança (TRV), apresentados na Tabela 6, foi complementada pela inspeção direta dos componentes estimados da matriz de variância-covariância dos efeitos aleatórios, detalhados no Apêndice B.

Inicialmente, observa-se que a estrutura identidade (pdIdent) produziu estimativas praticamente nulas para as variâncias dos efeitos aleatórios nos três modelos avaliados. No modelo de Ørskov, por exemplo, obteve-se $\text{Var}(\beta_1) = \text{Var}(\beta_2) = \text{Var}(\beta_3) = 0,000441$, valores incompatíveis com a magnitude da heterogeneidade observada entre curvas individuais. Comportamento semelhante foi verificado nos modelos de Gompertz e Logístico, com variâncias igualmente próximas de zero. Sob essa estrutura, a variância residual permaneceu elevada (30,79; 73,75 e 98,18, respectivamente), indicando que a imposição de variâncias homogêneas levou à absorção indevida da variabilidade estrutural pelo termo de erro.

A adoção da estrutura diagonal (pdDiag) revelou heterogeneidade substancial entre os parâmetros. No modelo de Ørskov, observou-se $\text{Var}(\beta_2) = 197,61$, contrastando com $\text{Var}(\beta_3) = 0,000617$, evidenciando que a maior parcela da variabilidade entre unidades experimentais concentra-se no parâmetro associado ao potencial de produção. Padrão análogo foi identificado nos modelos de Gompertz e Logístico, nos quais a variância de β_2 superou amplamente as demais. Além disso, a variância residual foi substancialmente reduzida sob pdDiag, indicando decomposição mais adequada da variabilidade total.

A estrutura pdSymm, por sua vez, permitiu estimar explicitamente as correlações entre os efeitos aleatórios associados aos parâmetros do modelo. No modelo de Ørskov, as correlações estimadas foram moderadas ($\text{Corr}(\beta_1, \beta_2) = -0,366$, $\text{Corr}(\beta_1, \beta_3) = 0,082$, $\text{Corr}(\beta_2, \beta_3) = -0,477$). Para os modelos de Gompertz e Logístico, observaram-se algumas correlações de elevada magnitude, como $\text{Corr}(\beta_1, \beta_3) = 0,954$ e $\text{Corr}(\beta_1, \beta_2) = 0,988$, indicando possível dependência entre parâmetros e sugerindo cautela na adoção de estruturas mais complexas para a matriz **D**.

Embora a estrutura **pdSymm** tenha produzido pequenos ganhos de log-verossimilhança e reduções marginais nos critérios de informação em alguns casos, particularmente no modelo de Ørskov, tais melhorias foram numericamente modestas quando comparadas ao aumento no número de parâmetros estimados. A diferença entre **pdDiag** e **pdSymm** em termos de AIC e BIC mostrou-se reduzida, sugerindo benefício estatístico limitado da inclusão das covariâncias.

Os testes da razão de verossimilhança indicaram que a transição da estrutura **pdIdent** para **pdDiag** promoveu melhoria altamente significativa no ajuste dos modelos (todos com $p < 0,0001$). Por outro lado, a comparação entre **pdDiag** e **pdSymm** não evidenciou melhora estatisticamente significativa para os modelos de Gompertz ($p = 0,2893$) e Logístico ($p = 0,0872$), sendo marginal apenas no modelo de Ørskov ($p = 0,0073$).

Dessa forma, considerando conjuntamente os critérios de informação, a magnitude das variâncias estimadas, a presença de correlações elevadas em algumas estruturas e o princípio da parcimônia, optou-se pela adoção da estrutura diagonal (**pdDiag**) como especificação final da matriz **D** para os três modelos avaliados.

No caso do modelo de Ørskov & McDonald, a adoção da estrutura diagonal para a matriz de variância-covariância dos efeitos aleatórios (**D**) também teve como objetivo manter padronização estrutural entre as funções avaliadas, assegurando comparabilidade direta entre os modelos. A especificação alternativa com matriz geral positiva definida (**pdSymm**) não apresentou ganho substancial em termos de ajuste global ou desempenho preditivo. Adicionalmente, para fins comparativos, procedeu-se ao ajuste do modelo com a inclusão explícita de uma estrutura de correlação residual contínua do tipo gaussiana, associada a uma variância residual constante (estrutura identidade) na matriz **R**. Essa especificação evidenciou instabilidade numérica quando combinada a estruturas mais complexas para **D** (Apêndice D), reforçando a decisão de adotar uma estrutura mais parcimoniosa.

Tabela 6 – Comparação das estruturas da matriz de variância-covariância dos efeitos aleatórios (**D**) para os modelos com efeitos aleatórios em β_1 , β_2 e β_3 .

Modelo	d	Estrutura	df	AIC	BIC	logLik	Teste	TRV	valor p
OR _{7,d}	1	pdIdent	35	4722,870	4882,900	-2326,435	—	—	—
	2	pdDiag	37	4255,361	4424,535	-2090,680	1 vs 2	471,509	<0,0001
	3	pdSymm	40	4249,338	4432,229	-2084,669	2 vs 3	12,023	0,0073
GO _{7,d}	1	pdIdent	35	5314,851	5474,881	-2622,425	—	—	—
	2	pdDiag	37	5123,165	5292,339	-2524,582	1 vs 2	195,686	<0,0001
	3	pdSymm	40	5125,411	5308,303	-2522,706	2 vs 3	3,754	0,2893
LO _{7,d}	1	pdIdent	35	5512,183	5672,213	-2721,091	—	—	—
	2	pdDiag	37	5358,440	5527,614	-2642,220	1 vs 2	157,743	<0,0001
	3	pdSymm	40	5357,875	5540,766	-2638,938	2 vs 3	6,564	0,0872

Nota: df = graus de liberdade; AIC = Critério de Informação de Akaike; BIC = Critério de Informação Bayesiano; logLik = log-verossimilhança; TRV = teste da razão de verossimilhança entre modelos aninhados. **pdIdent** = matriz identidade; **pdDiag** = matriz diagonal; **pdSymm** = matriz geral positiva-definida. Valores de $p < 0,05$ indicam melhora estatisticamente significativa.

A análise dos resíduos padronizados em função dos valores ajustados indicou que os três modelos apresentaram comportamentos residuais indesejáveis. No modelo $OR_{7,2}$, observou-se a presença de padrões estruturados ao longo dos valores ajustados, com alternância de regiões de resíduos positivos e negativos, sugerindo que a variabilidade não é totalmente explicada pela componente média do modelo. No modelo $GO_{7,2}$, esse comportamento mostrou-se mais pronunciado, com tendência sistemática e maior dispersão dos resíduos em determinados intervalos, indicando possível heterocedasticidade e inadequação local do ajuste. De forma semelhante, o modelo $LO_{7,2}$ também apresentou estrutura residual persistente, caracterizada por agrupamentos e desvios sistemáticos em relação à linha de referência zero.

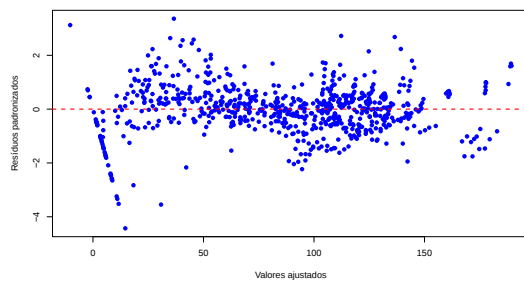
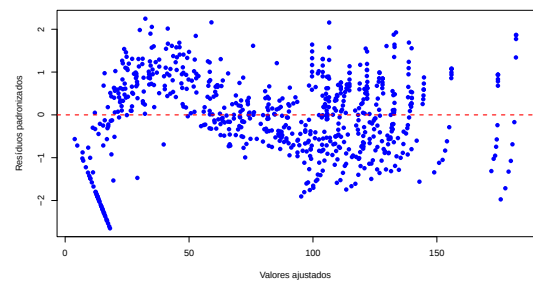
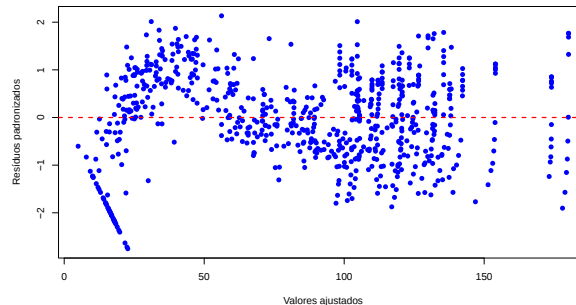
(a) Modelo de Ørskov & McDonald (1979) ($OR_{7,2}$)(b) Modelo de Gompertz ($GO_{7,2}$)(c) Modelo Logístico ($LO_{7,2}$)

Figura 9 – Gráficos de resíduos padronizados em função dos valores ajustados dos MNLM

Em todos os casos, a ausência de uma distribuição aproximadamente aleatória dos resíduos em torno de zero sugere que, apesar da incorporação de efeitos aleatórios, os modelos não lineares mistos avaliados ainda não descrevem plenamente a dinâmica dos dados ao longo do domínio dos valores ajustados. Os gráficos de autocorrelação (ACF) dos resíduos (Figura 10) confirmam a presença de autocorrelação serial significativa em todos os modelos ajustados. É possível observar que diversos coeficientes de autocorrelação ultrapassam os limites de significância, especialmente para defasagens iniciais, evidenciando a necessidade de incluir estruturas de correlação nos resíduos. Essa constatação reforça os indícios já observados nos gráficos de resíduos em função do tempo, apontando para a violação do pressuposto de independência dos erros.

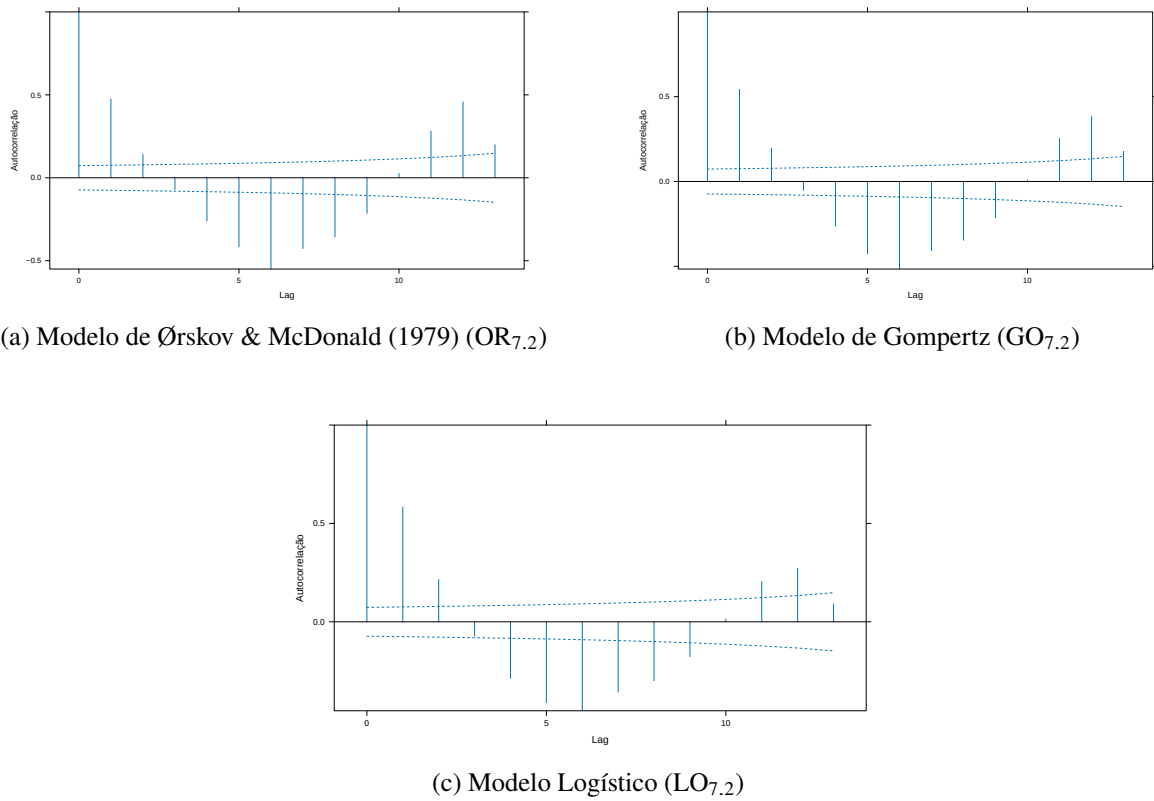


Figura 10 – Autocorrelação dos resíduos dos modelos não lineares mistos por estrutura de tempo

A inspeção da função de autocorrelação dos resíduos condicionais, realizada previamente à especificação de estruturas adicionais para a matriz de covariância residual $\mathbf{R}$, indicou a presença de dependência temporal intraunidade, ainda que sem um padrão estritamente monotônico nos defasamentos avaliados. Tal evidência motivou a avaliação formal de diferentes especificações para a matriz residual.

Foram consideradas, além da estrutura independente $\sigma^2\mathbf{I}$, diferentes estruturas de correlação, incluindo a simetria composta (CS) e funções contínuas de correlação, representadas pelas estruturas exponencial (EXP) e gaussiana (GAUS). Adicionalmente, avaliou-se a incorporação de heterogeneidade de variâncias entre tratamentos por meio da função `varIdent()` (PINHEIRO; BATES; R Core Team, 2025). Os critérios comparativos encontram-se apresentados na Tabela 7, enquanto os componentes estimados da matriz residual são detalhados no Apêndice C.

Observa-se que, para os três modelos não lineares avaliados, as estruturas contínuas de correlação (EXP e GAUS) promoveram redução substancial dos valores de AIC e aumento expressivo da log-verossimilhança quando comparadas à especificação independente. No modelo de Ørskov & McDonald, por exemplo, o AIC reduziu de 4255,361 para 3597,929 com a adoção da estrutura exponencial, evidenciando melhora significativa no ajuste. Comportamento análogo foi verificado nos modelos de Gompertz e Logístico, nos quais as reduções de AIC também foram expressivas, reforçando a presença de dependência temporal relevante nos dados.

A inclusão de heterogeneidade de variâncias por meio da estrutura `varIdent()` resultou em melhorias adicionais nos critérios de informação para ambas as funções de correlação contínua. Em particular, observa-se que os modelos com EXP + ID e GAUS + ID apresentaram valores idênticos de AIC, BIC e log-verossimilhança no modelo de Ørskov & McDonald, sugerindo que a principal contribuição para o ganho de ajuste decorre da modelagem da heterogeneidade de variâncias, enquanto a escolha entre as funções de correlação contínua exerce impacto secundário.

Em contraste, a estrutura de simetria composta (CS) não apresentou melhora relevante nos critérios de informação quando comparada à especificação independente, indicando ausência de evidência empírica para a hipótese de correlação constante ao longo do tempo. Ainda que a combinação CS + ID tenha proporcionado alguma redução nos critérios de informação, tal estrutura não apresenta suporte teórico consistente para dados longitudinais com dependência temporal decrescente, como é o caso das medições de fermentação *in vitro*.

Tabela 7 – Critérios de seleção dos modelos não lineares mistos com diferentes estruturas da matriz de covariância dos erros (**R**).

Modelo	r	Matriz R	df	AIC	BIC	logLik
OR _{7.2.r}	0	$\sigma^2\mathbf{I}$	37	4255,361	4424,535	-2090,680
	1	CS	38	4257,361	4431,107	-2090,680
	2	CS + ID	53	4102,183	4344,514	-1998,092
	3	EXP	38	3597,929	3771,676	-1760,965
	4	EXP + ID	53	3519,139	3761,470	-1706,569
	5	GAUS	38	3691,990	3865,737	-1807,995
	6	GAUS + ID	53	3519,139	3761,470	-1706,569
GO _{7.2.r}	0	$\sigma^2\mathbf{I}$	37	5123,165	5292,339	-2524,582
	1	CS	38	5125,165	5298,911	-2524,582
	2	CS + ID	53	5075,960	5318,291	-2484,980
	3	EXP	38	4487,931	4661,678	-2205,966
	4	EXP + ID	53	4462,936	4705,267	-2178,468
	5	GAUS	38	4440,370	4614,117	-2182,185
	6	GAUS + ID	53	4407,494	4649,825	-2150,747
LO _{7.2.r}	0	$\sigma^2\mathbf{I}$	37	5358,440	5527,614	-2642,220
	1	CS	38	5360,440	5534,186	-2642,220
	2	CS + ID	53	5330,147	5572,478	-2612,073
	3	EXP	38	4704,329	4878,076	-2314,165
	4	EXP + ID	53	4663,657	4905,988	-2278,828
	5	GAUS	38	4528,015	4701,762	-2226,008
	6	GAUS + ID	53	4509,525	4751,856	-2201,762

Nota: df = graus de liberdade; AIC = Critério de Informação de Akaike; BIC = Critério de Informação Bayesiano; logLik = log-verossimilhança. CS = simetria composta; EXP = Exponencial, GAUS = Guassiana, ID = variância heterogênea; estruturas Exponencial e Guassiana correspondem a funções contínuas de correlação.

Conforme discutido por Pinheiro e Bates (2000), a seleção da estrutura residual deve

considerar não apenas a parcimônia, mas também a adequação ao mecanismo de dependência presente nos dados. Diante da expressiva melhora nos critérios de informação, da coerência com a natureza temporal das observações e da capacidade de modelar simultaneamente a dependência e a heterogeneidade de variâncias, optou-se pela adoção da estrutura gaussiana com variâncias heterogêneas (GAUS + ID) como especificação final para **R** nos três modelos não lineares mistos (OR_{7.2.6}, GO_{7.2.6} e LO_{7.2.6}).

Nos gráficos half-normal com envelopes simulados (HNP) dos resíduos condicionais padronizados para os modelos de Ørskov & McDonald, Gompertz e Logístico (Figura 11) observa-se que o modelo de Ørskov & McDonald apresenta comportamento residual adequado, com a maior parte dos pontos contida dentro das bandas de confiança, indicando bom ajuste do modelo aos dados.

Para o modelo de Gompertz, nota-se que diversos pontos encontram-se próximos às bordas dos envelopes simulados, sugerindo um ajuste limitado, ainda que sem evidência clara de violação severa das suposições do modelo. Esse comportamento indica a presença de pequenas discrepâncias entre os valores observados e os esperados sob o modelo ajustado.

Já o modelo Logístico apresenta maior quantidade de pontos fora das bandas de confiança, especialmente nas regiões extremas, evidenciando inadequação da estrutura residual assumida. Esse padrão sugere que o modelo não captura de forma satisfatória a variabilidade presente nos dados, resultando em resíduos com comportamento não compatível com a distribuição teórica esperada.

Sob a perspectiva diagnóstica global, os resultados indicam superioridade do modelo de Ørskov & McDonald em relação aos demais, seguido pelo modelo de Gompertz, enquanto o modelo Logístico apresenta desempenho inferior no que se refere ao ajuste residual.

Adicionalmente, em análises preliminares conduzidas sem a especificação de uma estrutura de covariância residual, foram observadas instabilidades relevantes nas estimativas dos parâmetros, incluindo valores negativos para componentes biologicamente não negativos, como a fração solúvel (β_1).

A incorporação explícita de estruturas de correlação e heterogeneidade de variâncias permitiu modelar a dependência temporal inerente às observações longitudinais, resultando em maior estabilidade numérica, eliminação de estimativas biologicamente incoerentes e maior coerência inferencial dos modelos ajustados. Dessa forma, a continuidade da análise fundamentou-se nos modelos com estrutura residual especificada, considerados estatisticamente mais adequados.

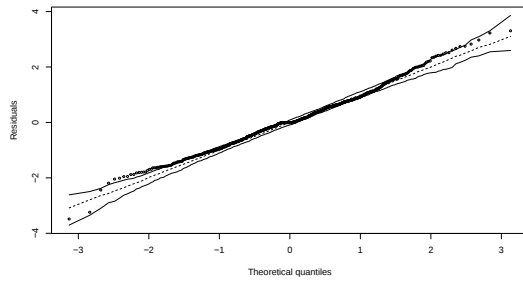
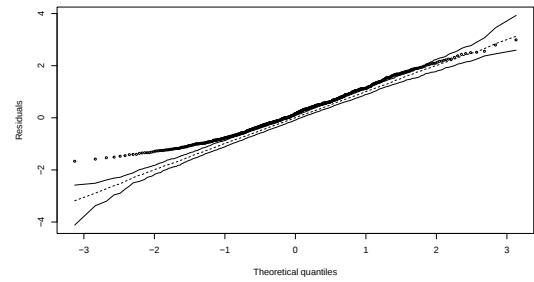
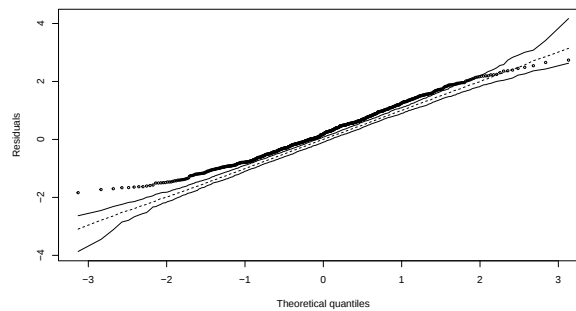
(a) Modelo de Ørskov & McDonald (OR_{7.2.6})(b) Modelo de Gompertz (GO_{7.2.6})(c) Modelo Logístico (LO_{7.2.6})

Figura 11 – Gráficos HNP dos resíduos dos modelos não lineares mistos ajustados.

A Figura 12 apresenta os gráficos de resíduos padronizados em função dos valores ajustados para os modelos com estrutura residual especificada. Em relação à Figura 9, na qual os ajustes consideravam apenas modificações na matriz de variância e covariância dos efeitos aleatórios (**D**), observa-se melhora geral no comportamento diagnóstico dos resíduos.

Resultados similares aos observados nos gráficos de envelope (Figura 11) também são verificados nesta análise complementar, na qual o modelo de Ørskov & McDonald mantém comportamento residual relativamente mais favorável entre os modelos avaliados. O modelo de Gompertz apresenta desempenho intermediário, enquanto o modelo Logístico evidencia maior dispersão residual e padrões menos compatíveis com os pressupostos adotados. Dessa forma, os gráficos de resíduos versus valores ajustados atuam como evidência adicional aos envelopes simulados, contribuindo para uma interpretação diagnóstica mais clara e consistente.

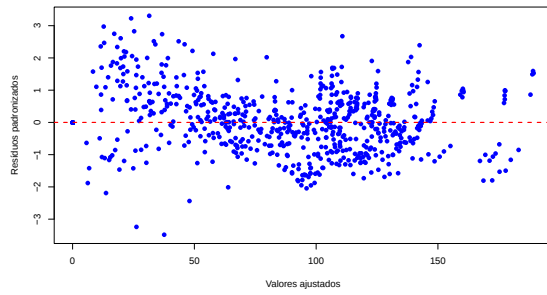
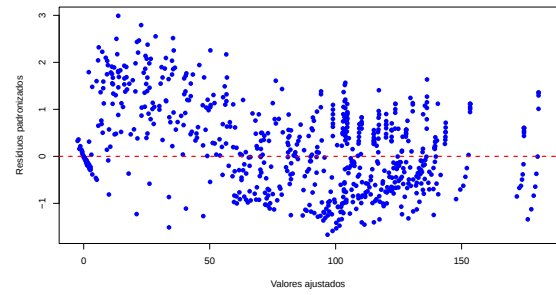
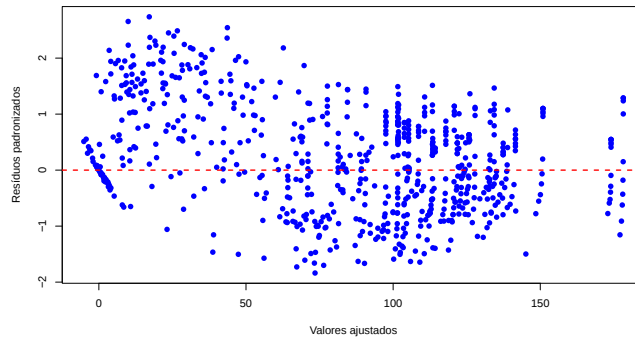
(a) Modelo de Ørskov & McDonald (OR_{7.2.6})(b) Modelo de Gompertz (GO_{7.2.6})(c) Modelo Logístico (LO_{7.2.6})

Figura 12 – Gráficos de resíduos padronizados versus valores ajustados dos modelos com estrutura gaussiana e variâncias heterogêneas.

Após a etapa de diagnóstico dos resíduos, procedeu-se à avaliação da capacidade preditiva dos modelos ajustados, considerando tanto a predição marginal quanto a condicional. Para isso, foi construída uma rotina computacional que permitiu gerar predições a partir dos modelos não lineares mistos e calcular métricas de desempenho associadas.

As predições foram obtidas em dois níveis: (i) nível marginal (level = 0), considerando apenas os efeitos fixos, e (ii) nível condicional (level = 1), incorporando simultaneamente os efeitos fixos e aleatórios. Essa distinção permite avaliar, respectivamente, a capacidade de generalização do modelo e sua aderência aos dados observados dentro das unidades experimentais.

A avaliação preditiva foi realizada com base em um subconjunto de dados de teste, composto por tempos não utilizados diretamente no ajuste dos modelos, garantindo assim uma análise mais robusta da capacidade de extrapolação. Para cada modelo, foram calculadas as seguintes métricas: raiz do erro quadrático médio (RMSE), erro absoluto médio (MAE), erro percentual absoluto médio (MAPE), coeficiente de correlação de Pearson (Corr) e coeficiente de concordância de Lin (CCC).

As métricas foram obtidas em nível global, considerando todas as observações do conjunto de teste, de modo a permitir a comparação direta entre os modelos avaliados, conforme

apresentado na Tabela 8. De modo geral, observa-se que os modelos apresentaram melhor desempenho no nível condicional (level = 1), como esperado, uma vez que esse nível incorpora os efeitos aleatórios específicos de cada unidade experimental. Por outro lado, a avaliação no nível marginal (level = 0) mostrou-se fundamental para a comparação entre modelos em termos de capacidade preditiva geral.

A Tabela 8 evidencia que o modelo de Ørskov & McDonald (OR_{7.2.6}) apresentou desempenho superior em todas as métricas consideradas, tanto no nível marginal quanto no condicional, com menores valores de RMSE, MAE e MAPE, e maiores coeficientes de correlação e concordância. O modelo de Gompertz (GO_{7.2.6}) apresentou desempenho intermediário, enquanto o modelo Logístico (LO_{7.2.6}) apresentou os maiores erros preditivos e menores níveis de concordância.

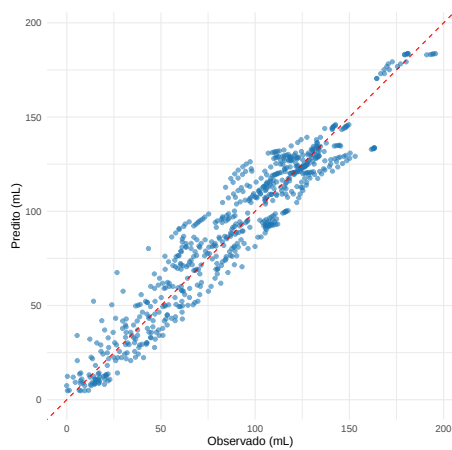
Resultados consistentes são observados na Figura 13, que apresenta os gráficos de valores observados versus preditos obtidos a partir do conjunto de dados de teste, considerando predições no nível marginal. Observa-se que, embora os três modelos acompanhem a tendência geral dos dados, o modelo de Ørskov apresenta distribuição mais próxima da linha de identidade, com menor dispersão vertical, especialmente na faixa intermediária dos valores.

Nos maiores níveis de produção de gás, nota-se leve compressão das predições, padrão esperado em modelos não lineares com comportamento assintótico. Ainda assim, o modelo OR_{7.2.6} mantém melhor aproximação global, corroborando as evidências numéricas apresentadas, que em conjunto com a inspeção gráfica reforça sua superioridade em relação aos outros modelos, tanto em termos de capacidade de generalização quanto de aderência aos dados observados. Dessa forma, essa especificação foi adotada como modelo de referência nas análises subsequentes.

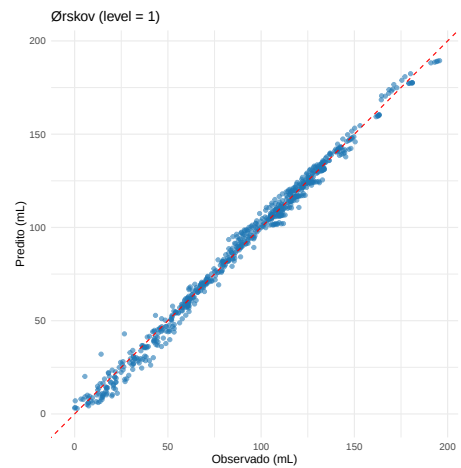
Tabela 8 – Métricas preditivas globais dos modelos não lineares mistos ajustados, avaliadas em conjunto de dados de teste.

Nível	Métrica	OR _{7.2.6}	GO _{7.2.6}	LO _{7.2.6}
0	RMSE	12,239	13,927	14,858
	MAE	9,734	11,284	12,091
	MAPE (%)	24,676	22,651	24,383
	Corr	0,958	0,951	0,947
	CCC	0,958	0,949	0,942
1	RMSE	3,924	8,629	10,936
	MAE	2,953	7,025	8,967
	MAPE (%)	10,817	16,622	20,472
	Corr	0,996	0,985	0,975
	CCC	0,996	0,981	0,969

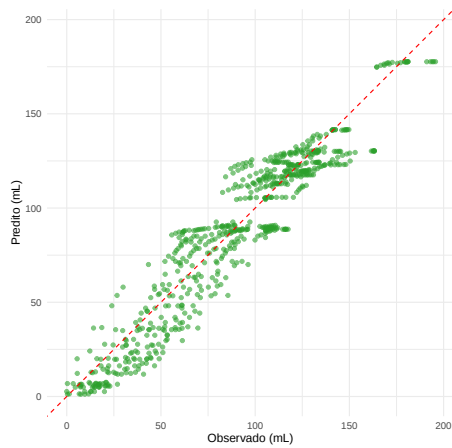
Nota: RMSE = raiz do erro quadrático médio; MAE = erro absoluto médio; MAPE = erro percentual absoluto médio; Corr = coeficiente de correlação de Pearson; CCC = coeficiente de concordância de Lin. Nível 0: predição marginal; Nível 1: predição condicional.



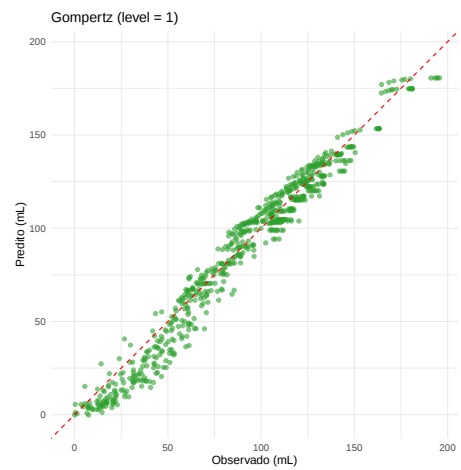
(a) Ørskov (level = 0)



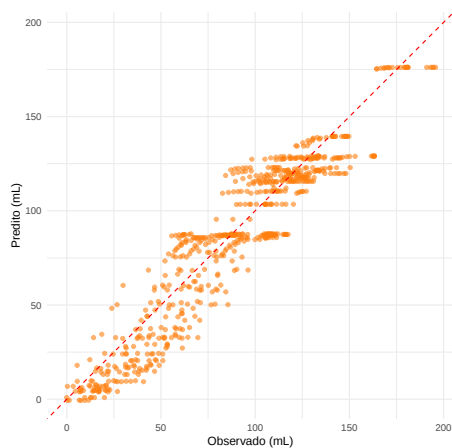
(b) Ørskov (level = 1)



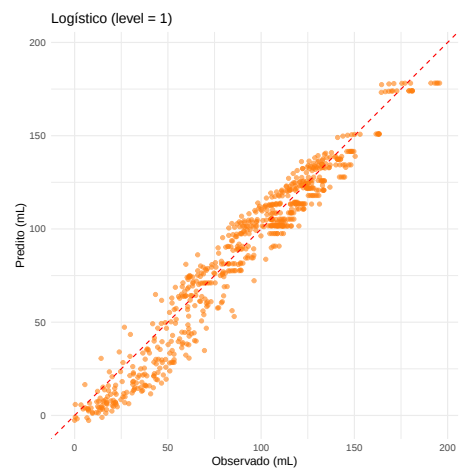
(c) Gompertz (level = 0)



(d) Gompertz (level = 1)



(e) Logístico (level = 0)



(f) Logístico (level = 1)

Figura 13 – Gráficos de valores preditos *versus* observados para os modelos não lineares mistos, considerando predições no nível marginal (level = 0) e condicional (level = 1), avaliados em conjunto de dados de teste.

A Tabela 9 apresenta as estimativas dos parâmetros do modelo de Ørskov & McDonald com a estrutura residual selecionada. Observa-se que o parâmetro associado à fração solúvel (β_1) apresenta baixa magnitude, acompanhada de erro-padrão relativamente elevado, sugerindo contribuição mais homogênea entre os tratamentos. Em contraste, o parâmetro referente à fração potencialmente degradável (β_2) apresenta variação expressiva, constituindo o principal componente responsável pelas diferenças no potencial total de degradação, com destaque para o tratamento 9, que apresentou a maior estimativa. O parâmetro β_3 , associado à taxa de degradação, apresentou variação moderada, com influência secundária na distinção entre tratamentos.

As estimativas de degradabilidade efetiva (DE), calculadas considerando taxa de passagem ruminal $k = 0,05$, reforçam esse padrão, com maiores valores associados a maiores estimativas de β_2 . A incorporação dos erros-padrão e intervalos de confiança permite avaliar a precisão das estimativas e a sobreposição entre tratamentos. Embora exista hierarquização em termos de valores pontuais, observa-se sobreposição parcial em alguns casos, indicando que nem todas as diferenças são estatisticamente pronunciadas. Ainda assim, o tratamento 9 destacou-se de forma consistente, apresentando o maior valor de DE.

As comparações múltiplas apresentadas no Apêndice E corroboram essas evidências. Para o parâmetro β_1 , não foram realizadas comparações entre tratamentos, uma vez que este parâmetro foi estimado como constante no modelo. Para β_3 , não foram identificadas diferenças estatisticamente significativas entre os tratamentos pelo teste de Tukey ($p > 0,05$), indicando comportamento relativamente uniforme da taxa de degradação entre as amostras avaliadas.

Tabela 9 – Estimativas dos parâmetros do modelo de Ørskov e McDonald (OR_{7.2.6}) e degradabilidade efetiva (DE).

T_i	$\hat{\beta}_1$		$\hat{\beta}_2$		$\hat{\beta}_3$		$\hat{\beta}_1 + \hat{\beta}_2$		DE		IC 95% DE	
	Est	EP	Est	EP	Est	EP	Est	EP	Est	EP	Inf	Sup
1 ⁽¹³⁾	0,017	0,466	101,4	8,32	0,0762	0,0141	101,46	8,32	61,27	6,75	48,03	74,51
2 ⁽¹⁴⁾	0,017	0,466	93,3	8,00	0,0911	0,0164	93,35	7,99	60,27	6,45	47,63	72,92
3 ⁽⁸⁾	0,017	0,466	121,3	9,41	0,0927	0,0181	121,27	9,40	78,78	8,16	62,78	94,78
4 ⁽¹⁵⁾	0,017	0,466	103,2	7,87	0,0493	0,0138	103,26	7,87	51,27	8,21	35,18	67,37
5 ⁽⁶⁾	0,017	0,466	132,8	7,67	0,0861	0,0141	132,82	7,66	84,04	7,01	70,29	97,78
6 ⁽³⁾	0,017	0,466	135,5	7,65	0,1022	0,0143	135,55	7,64	91,01	6,64	78,00	104,03
7 ⁽⁵⁾	0,017	0,466	134,3	7,63	0,0935	0,0141	134,27	7,62	87,48	6,80	74,16	100,80
8 ⁽¹⁰⁾	0,017	0,466	131,5	7,65	0,0679	0,0138	131,50	7,64	75,75	7,90	60,27	91,23
9⁽¹⁾	0,017	0,466	184,1	9,26	0,1041	0,0170	184,14	9,25	124,38	9,09	106,56	142,20
10 ⁽⁹⁾	0,017	0,466	109,2	9,31	0,1205	0,0182	109,22	9,30	77,20	7,43	62,64	91,75
11 ⁽¹²⁾	0,017	0,466	121,5	7,83	0,0790	0,0142	121,51	7,82	74,41	7,08	60,54	88,29
12 ⁽⁷⁾	0,017	0,466	130,6	7,65	0,0795	0,0140	130,64	7,64	80,21	7,21	66,08	94,33
13 ⁽¹⁶⁾	0,017	0,466	101,9	7,83	0,0479	0,0138	101,95	7,83	49,91	8,28	33,68	66,14
14 ⁽²⁾	0,017	0,466	145,8	10,30	0,0888	0,0179	145,80	10,31	93,27	9,47	74,70	111,83
15 ⁽⁴⁾	0,017	0,466	147,4	9,25	0,0737	0,0168	147,37	9,25	87,83	9,82	68,59	107,06
16 ⁽¹¹⁾	0,017	0,466	125,7	9,21	0,0735	0,0168	125,73	9,21	74,83	8,84	57,50	92,16

Est = estimativa; EP = erro-padrão. A degradabilidade efetiva (DE) foi calculada considerando taxa de passagem ruminal $k = 0,05$. Índice sobrescrito indica a posição relativa baseada na DE (1 = maior valor).

Por outro lado, para o parâmetro β_2 foram observadas diferenças estatisticamente significativas entre tratamentos, evidenciando que a principal fonte de variação no processo fermentativo está associada à fração potencialmente degradável. Em particular, o tratamento 9 apresentou contrastes significativos em relação a diversos tratamentos, confirmando seu maior potencial de degradação e sua posição de destaque na hierarquização dos valores de degradabilidade efetiva.

Esses resultados reforçam a interpretação de que a discriminação entre tratamentos ocorre predominantemente em função da magnitude de β_2 , enquanto β_1 e β_3 desempenham papel mais uniforme no conjunto de dados analisado. A Figura 14 sintetiza esses resultados ao apresentar as curvas ajustadas por tratamento e réplica com base no modelo final, nas quais tratamentos com maiores valores de β_2 e DE apresentam maior acúmulo de gás ao longo do tempo de incubação.

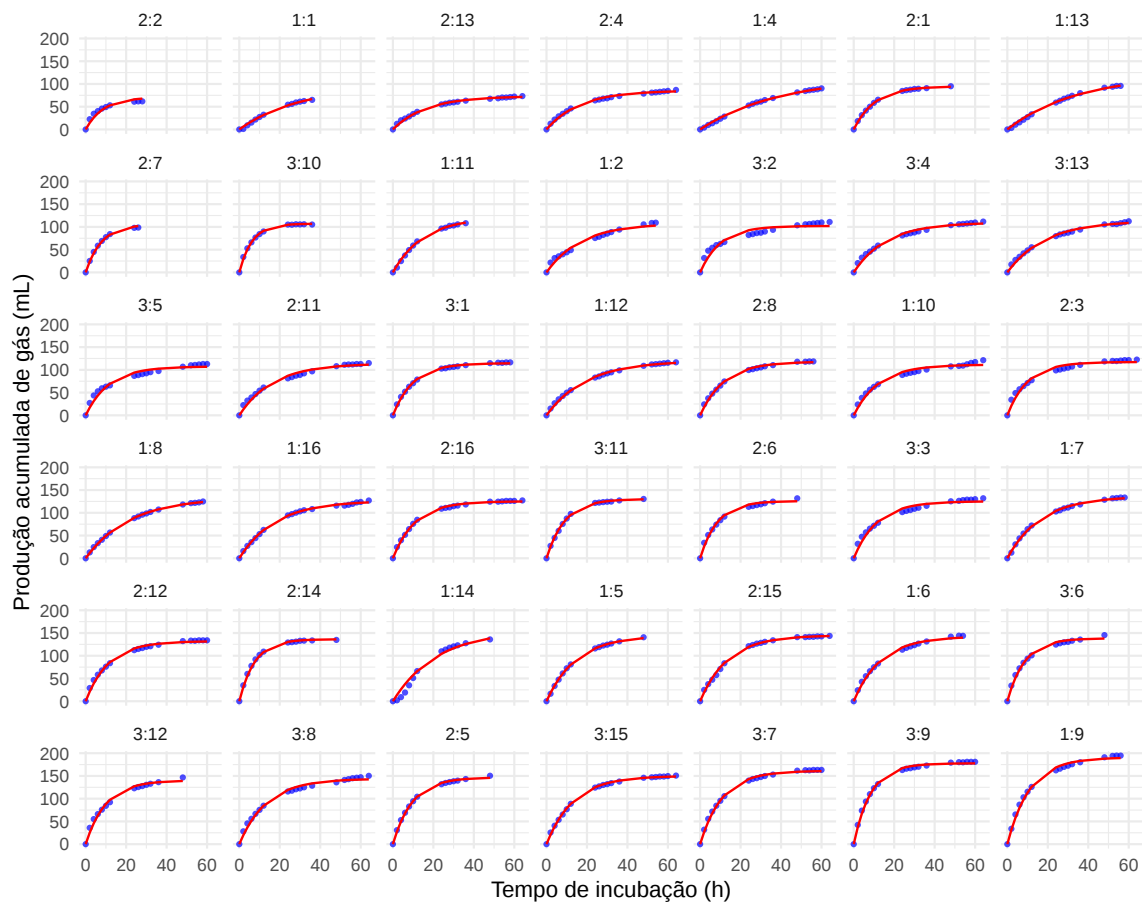


Figura 14 – Valores observados e curvas ajustadas por réplica e tratamento segundo o modelo OR_{7.2.6}.

CONCLUSÃO

Os modelos não lineares mistos mostraram-se adequados para a análise da cinética de produção de gás *in vitro*, permitindo descrever de forma consistente o comportamento fermentativo das amostras ao longo do tempo e capturar a variabilidade inerente aos dados longitudinais. A modelagem hierárquica possibilitou representar simultaneamente a estrutura média da curva e a heterogeneidade entre unidades experimentais, resultando em estimativas estáveis e biologicamente interpretáveis.

Entre os modelos avaliados, o modelo de Ørskov & McDonald apresentou desempenho preditivo superior quando consideradas conjuntamente as métricas raiz do erro quadrático médio (RMSE), erro absoluto médio (MAE), erro percentual absoluto médio (MAPE), correlação de Pearson (Corr) e coeficiente de concordância de Lin (CCC) ao conjunto de dados teste. A superioridade desse modelo foi evidenciada tanto no nível marginal quanto no nível condicional de predição, destacando sua capacidade de representar o comportamento médio e a variabilidade individual das unidades experimentais.

A inclusão de efeitos aleatórios nos parâmetros cinéticos mostrou-se essencial para capturar a heterogeneidade entre amostras, particularmente nos parâmetros associados à fração potencialmente degradável e à taxa de degradação, componentes diretamente relacionados ao potencial fermentativo e à dinâmica de utilização do substrato. A estrutura diagonal para a matriz de variância-covariância dos efeitos aleatórios (**D**) revelou-se adequada e parcimoniosa, não havendo evidência empírica consistente que justificasse a adoção de estruturas mais complexas.

No que se refere à matriz residual (**R**), a modelagem da dependência intraunidade experimental e da heterogeneidade de variâncias mostrou-se fundamental para a obtenção de estimativas estáveis. A incorporação de estruturas de correlação associadas à função de variância permitiu melhor representação da variabilidade dos dados, contribuindo para maior coerência inferencial. A avaliação diagnóstica por meio de envelopes simulados half-normal (HNP) indicou adequação global dos modelos ajustados, reforçando a consistência das suposições adotadas.

Quanto à comparação entre tratamentos, observou-se que as principais diferenças estão associadas à fração potencialmente degradável, enquanto a fração solúvel apresentou compor-

tamento relativamente homogêneo entre as amostras. A integração dos parâmetros por meio da degradabilidade efetiva (DE), calculada considerando taxa de passagem ruminal $k = 0,05$, permitiu sintetizar o desempenho fermentativo global. Destacaram-se os tratamentos 9, 6 e 14, que apresentaram os maiores valores de DE, refletindo maior potencial fermentativo *in vitro*, enquanto os menores desempenhos foram observados nos tratamentos 13 e 4.

De forma geral, os resultados demonstram que a combinação do modelo de Ørskov & McDonald com efeitos aleatórios nos parâmetros cinéticos, estrutura diagonal para **D** e modelagem adequada da estrutura residual **R** constitui uma estratégia estatisticamente robusta, biologicamente coerente e metodologicamente consistente para a análise da produção de gás *in vitro*. Essa abordagem aprimora a precisão das estimativas, melhora a capacidade preditiva do modelo e fornece subsídios quantitativos confiáveis para a comparação de alimentos em estudos de nutrição de ruminantes.

Assim, o modelo selecionado não apenas descreve adequadamente a dinâmica fermentativa observada, mas também se consolida como ferramenta analítica consistente para apoiar decisões relacionadas à formulação de dietas, avaliação de ingredientes e otimização da eficiência ruminal em sistemas de produção animal.

CONSIDERAÇÕES FUTURAS

Os resultados obtidos nesta dissertação demonstram que modelos não lineares mistos constituem uma estrutura estatística robusta para a análise da cinética de produção de gás *in vitro*, ao permitirem a modelagem simultânea da trajetória média do processo fermentativo e da variabilidade inerente aos dados longitudinais. A incorporação de efeitos aleatórios e de estruturas adequadas para a matriz residual mostrou-se fundamental para a obtenção de estimativas numericamente estáveis e biologicamente coerentes, evidenciando a relevância de abordagens hierárquicas na modelagem de fenômenos fermentativos complexos.

Todavia, as dificuldades observadas durante o processo de ajuste, como instabilidade paramétrica sob determinadas especificações da matriz residual, elevada correlação entre componentes dos efeitos aleatórios e ocorrência de estimativas biologicamente implausíveis quando a dependência intraunidade não era adequadamente modelada, indicam que a estrutura de variância-covariância desempenha papel central na qualidade inferencial dos modelos. Nesse sentido, uma agenda de pesquisa futura pode concentrar-se na investigação sistemática de diferentes parametrizações da matriz $\mathbf{R}$, incluindo estruturas que acomodem simultaneamente correlação temporal e heterocedasticidade. Destaca-se, em particular, a necessidade de explorar estruturas residuais capazes de reproduzir padrões não lineares observados nos diagnósticos gráficos, com comportamento sistemático evidenciado nos gráficos de resíduos e nas funções de autocorrelação, indicando que a dependência temporal pode não ser adequadamente capturada por estruturas convencionais.

Além disso, a comparação entre abordagens frequentistas e bayesianas para estimação desses componentes pode ampliar a flexibilidade da modelagem e melhorar a representação da dependência residual em contextos mais complexos.

Uma linha promissora de aprofundamento consiste na ampliação da estrutura hierárquica do modelo, incorporando múltiplos níveis de aleatoriedade, tais como efeitos associados a lotes experimentais, períodos de incubação, fontes de inóculo ruminal ou variabilidade intra-repetição. A modelagem multinível pode permitir uma decomposição mais detalhada das fontes de variabilidade biológica e experimental, contribuindo para maior precisão na inferência e

ampliando o potencial de generalização dos resultados.

Outra perspectiva relevante envolve a extensão das funções de crescimento avaliadas para formulações mais flexíveis, como modelos bicompartimentais mistos, modelos com parâmetros dependentes de covariáveis nutricionais ou abordagens não lineares semiparamétricas. Essas estruturas podem capturar com maior fidelidade transições entre fases fermentativas distintas, possibilitando interpretações mais refinadas sobre a dinâmica de degradação das frações rapidamente fermentáveis e estruturalmente complexas.

Adicionalmente, a integração entre modelagem estatística e caracterização físico-química detalhada dos substratos, incluindo análises bromatológicas e perfil de fibra, pode estabelecer um arcabouço quantitativo mais abrangente para explicar diferenças nos parâmetros cinéticos. Tal integração abre espaço para a construção de modelos preditivos que relacionem diretamente propriedades químicas dos alimentos aos parâmetros fermentativos estimados.

Por fim, a incorporação de abordagens bayesianas hierárquicas e técnicas de inferência computacional intensiva, como métodos baseados em amostragem MCMC, constitui um caminho promissor para lidar com estruturas de dependência mais complexas. Essa transição metodológica pode ampliar a flexibilidade da modelagem e permitir inferências probabilísticas completas sobre os parâmetros cinéticos e suas incertezas.

Dessa forma, o presente trabalho estabelece não apenas uma avaliação comparativa entre modelos não lineares mistos aplicados à produção de gás *in vitro*, mas também um ponto de partida para o desenvolvimento de estruturas hierárquicas mais abrangentes, capazes de integrar múltiplas fontes de variabilidade biológica e experimental, consolidando a modelagem estatística como ferramenta estratégica na pesquisa em nutrição de ruminantes.

REFERÊNCIAS

- ANDRADE, D. F.; SINGER, J. M. Análise de dados longitudinais. *Revista de Matemática e Estatística*, v. 4, n. 1, p. 1–23, 1986.
- ARCHONTOULIS, S. V.; MIGUEZ, F. E. Nonlinear regression models and applications in agricultural research. *Agronomy Journal*, Wiley Online Library, v. 107, n. 2, p. 786–798, 2015.
- BROCKWELL, P. J.; DAVIS, R. A. *Introduction to Time Series and Forecasting*. [S.l.]: Springer Science & Business Media, 2002.
- BURNHAM, K. P.; ANDERSON, D. R. *Model selection and multimodel inference: A practical information-theoretic approach*. 2. ed. New York: Springer Science & Business Media, 2002. ISBN 978-0-387-95364-9.
- CABRAL, Í. d. S. et al. Avaliação de modelos utilizados na produção de gás in vitro de alimentos tropicais. 2019.
- CASELLA, G.; BERGER, R. L. *Statistical Inference*. [S.l.]: Duxbury Press, 2002.
- DAVIDIAN, M.; GILTINAN, D. M. *Nonlinear Models for Repeated Measurement Data*. Boca Raton: Chapman & Hall/CRC, 2003.
- DEMIDENKO, E. *Mixed models: theory and applications with R*. [S.l.]: John Wiley & Sons, 2013.
- FAO. *The future of food and agriculture: Trends and challenges*. [S.l.]: Food and Agriculture Organization of the United Nations, 2017.
- FILHO, S. d. C. V. et al. *Nutrição de Ruminantes*. Viçosa, MG: Editora Suprema, 2006.
- FRANCE, J.; DIJKSTRA, J. *Quantitative Aspects of Ruminant Digestion and Metabolism*. Wallingford, UK: CABI Publishing, 2005.
- FRANCE, J.; THORNLEY, J. H. M. Mathematical models in agriculture. *Butterworths*, London, 1984.
- GU, X. et al. Describing and form digestibility in multiple pool models with focus on undf and modeling considerations for different feed types. *Animals*, MDPI, v. 14, n. 1, p. 169, 2024. Disponível em: <<https://www.mdpi.com/2076-2615/14/1/169>>.

- HASTIE, T.; TIBSHIRANI, R.; FRIEDMAN, J. *The Elements of Statistical Learning*. [S.l.]: Springer, 2009.
- HENDERSON, C. R. Estimation of genetic parameters. *Annals of Mathematical Statistics*, v. 21, n. 2, p. 309–310, 1950.
- JANEIRO, V. et al. Nonlinear mixed models applied to ruminal degradability studies. *International Journal of Psychological Studies*, Canadian Center of Science and Education, v. 11, n. 5, p. 1–18, 2022.
- KAMALAK, A. et al. Comparison of in vitro gas production technique with in situ nylon bag technique to estimate dry matter degradation. *Czech Journal of Animal Science*, v. 50, n. 2, 2005.
- LAIRD, N. M.; WARE, J. H. Random-effects models for longitudinal data. *Biometrics*, Wiley, v. 38, n. 4, p. 963–974, 1982.
- LAWRENCE, I.; LIN, K. A concordance correlation coefficient to evaluate reproducibility. *Biometrics*, JSTOR, p. 255–268, 1989.
- LI, S. et al. Applications of undf in ration modeling and formulation. *Animals*, MDPI, v. 13, n. 5, p. 947, 2023. Disponível em: <<https://www.mdpi.com/2076-2615/13/5/947>>.
- LINDSTROM, M. J.; BATES, D. M. Nonlinear mixed effects models for repeated measures data. *Biometrics*, JSTOR, p. 673–687, 1990.
- LÓPEZ, S. et al. A comparison of mathematical models to describe the disappearance curves obtained by the in vitro gas production technique. *Animal Feed Science and Technology*, v. 79, n. 2, p. 145–159, 1999.
- MEDEIROS, S. D. S. de et al. Modelos não lineares mistos em ensaios de degradabilidade ruminal in situ. *Ciência Animal Brasileira/Brazilian Animal Science*, v. 21, 2020.
- MERTENS, D. R. Rate and extent of digestion. *Proceedings of the Cornell Nutrition Conference*, p. 13–28, 1993.
- MERTENS, D. R. Kinetic models of fiber digestion. In: MURPHY, J. J. (Ed.). *Progress in Dairy Science*. Nottingham, UK: Nottingham University Press, 2016. p. 135–152.
- MERTENS, D. R.; ELY, L. O. A dynamic model of fiber digestion and passage in the ruminant for evaluating forage quality. *Journal of Animal Science*, v. 54, n. 4, p. 895–905, 1982.
- ØRSKOV, E. R.; MCDONALD, I. The estimation of protein degradability in the rumen from incubation measurements weighted according to rate of passage. *The Journal of Agricultural Science*, v. 92, n. 2, p. 499–503, 1979.
- PELL, A.; SCHOFIELD, P. Use of computer models to interpret gas production profiles for forage and concentrate feeds evaluated under in vitro conditions. *Animal Feed Science and Technology*, v. 64, n. 1, p. 77–89, 1997.
- PINHEIRO, J.; BATES, D.; R Core Team. *nlme: Linear and Nonlinear Mixed Effects Models*. [S.l.], 2025. R package version 3.1-168. Disponível em: <<https://CRAN.R-project.org/package=nlme>>.

PINHEIRO, J. C.; BATES, D. M. *Mixed-Effects Models in S and S-PLUS*. New York: Springer-Verlag, 2000.

R Core Team. *R: A Language and Environment for Statistical Computing*. Vienna, Austria, 2025. Disponível em: <<https://www.R-project.org/>>.

SARTORIO, S. D. *Modelos não lineares mistos em estudos de degradabilidade ruminal in situ*. Tese (Doutorado) — Universidade de São Paulo, 2013.

SAVIAN, T. V. *Modelagem da cinética de degradação ruminal de alimentos para ruminantes*. Tese (Tese de Doutorado) — Universidade Federal de Santa Maria, Santa Maria, 2008.

SCHOFIELD, P.; PELL, A. A comparison of gas production methods for evaluating the digestibility of forages. *Animal Feed Science and Technology*, v. 49, p. 57–76, 1994.

SELF, S. G.; LIANG, K.-Y. Asymptotic properties of maximum likelihood estimators and likelihood ratio tests under nonstandard conditions. *Journal of the American Statistical Association*, Taylor & Francis, v. 82, n. 398, p. 605–610, 1987.

TEIXEIRA, U. H. G. et al. Modelos matemáticos para estimação dos parâmetros da cinética de degradação ruminal de concentrados protéicos. *Revista Brasileira de Saúde e Produção Animal*, v. 17, n. 1, p. 73–85, 2016. Disponível em: <<https://www.scielo.br/j/rbspa/a/CCjJDnRTvSwfGZjHYZ7M5Qt/?lang=pt>>.

THEODOROU, M. K. et al. A simple gas production method using a pressure transducer to determine the fermentation kinetics of ruminant feeds. *Animal Feed Science and Technology*, v. 48, n. 3–4, p. 185–197, 1994.

THIAGO, L. R. L. S. *Avaliação da degradabilidade ruminal de alimentos para ruminantes*. Dissertação (Dissertação de Mestrado) — Universidade Federal de Viçosa, Viçosa, 1994.

TOMICH, T. et al. Adaptação de uma técnica "in vitro" para descrição da cinética de degradação ruminal da matéria seca de volumosos. 2006.

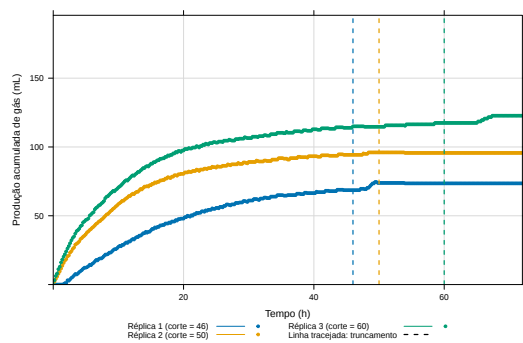
VERBEKE, G.; MOLENBERGHS, G. *Linear Mixed Models for Longitudinal Data*. [S.l.]: Springer Science & Business Media, 2009.

WEST, B. T.; WELCH, K. B.; GALECKI, A. T. *Linear Mixed Models: A Practical Guide Using Statistical Software*. [S.l.]: CRC Press, 2014.

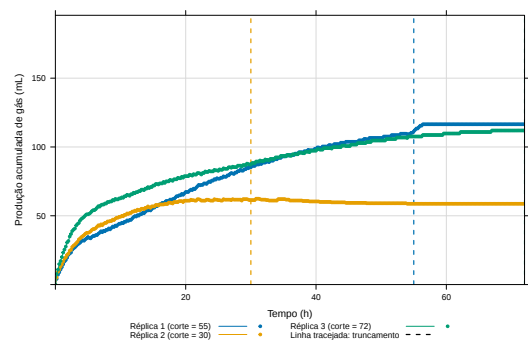
Apêndices

APÊNDICE A

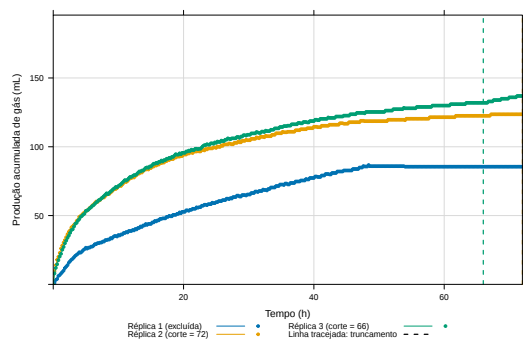
PROCEDIMENTO DE TRUNCAMENTO DAS CURVAS DE PRODUÇÃO DE GÁS



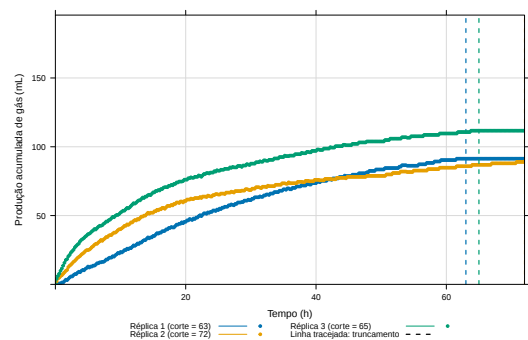
(a) Tratamento 1



(b) Tratamento 2

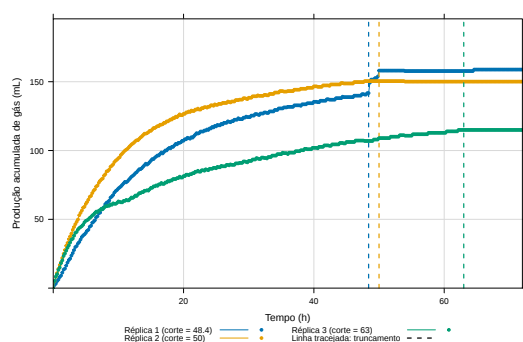


(c) Tratamento 3

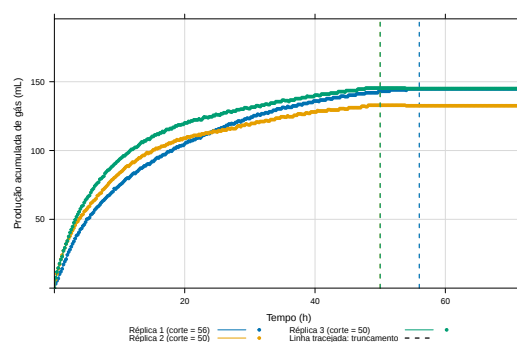


(d) Tratamento 4

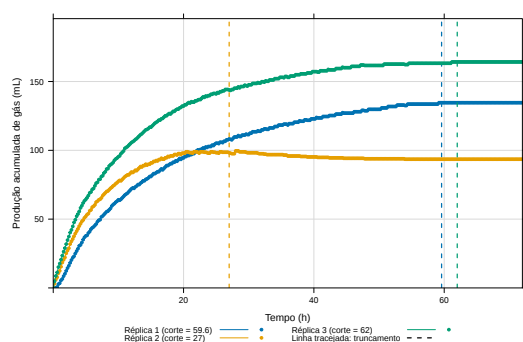
Figura 15 – Curvas de produção de gás e pontos de truncamento para os tratamentos avaliados de 1 a 4.



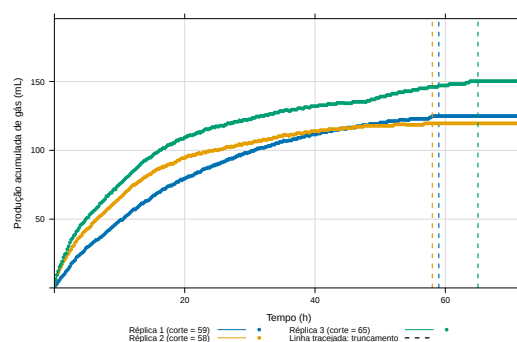
(a) Tratamento 5



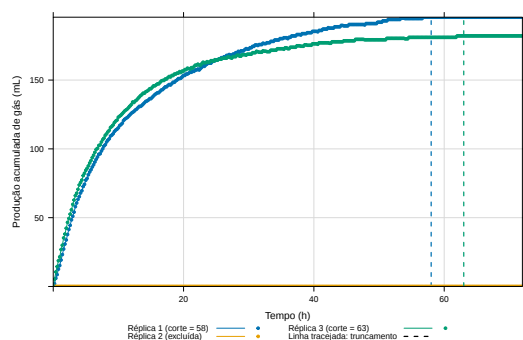
(b) Tratamento 6



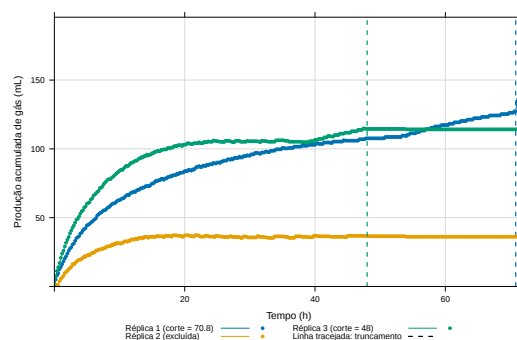
(c) Tratamento 7



(d) Tratamento 8

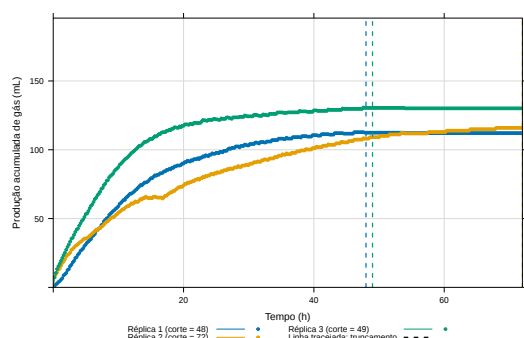


(e) Tratamento 9

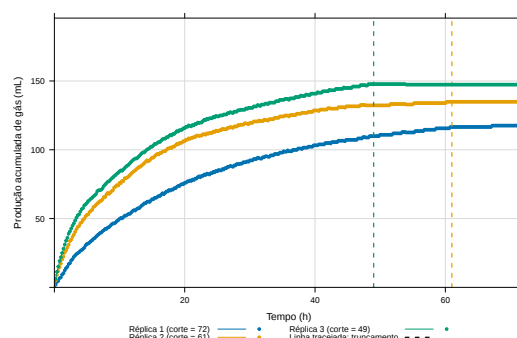


(f) Tratamento 10

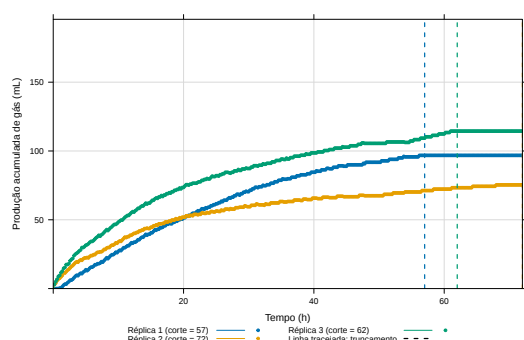
Figura 16 – Curvas de produção de gás e pontos de truncamento para os tratamentos avaliados de 5 a 10.



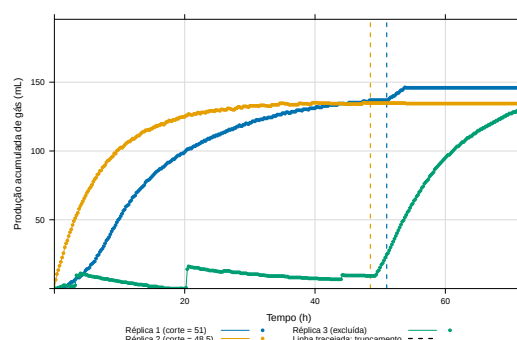
(a) Tratamento 11



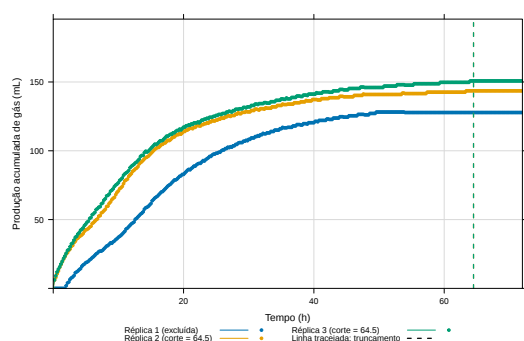
(b) Tratamento 12



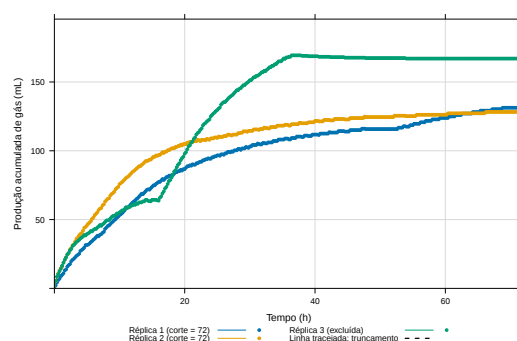
(c) Tratamento 13



(d) Tratamento 14



(e) Tratamento 15



(f) Tratamento 16

Figura 17 – Curvas de produção de gás e pontos de truncamento para os tratamentos avaliados de 11 a 16.

ESTRUTURAS DA MATRIZ DE VARIÂNCIA COVARIÂNCIA DOS EFEITOS ALEATÓRIOS

Tabela 10 – Componentes da matriz de variância–covariância dos efeitos aleatórios (**D**) e variância residual para os modelos não lineares mistos, considerando efeitos aleatórios em β_1 , β_2 e β_3 .

Modelo	Componente	pdIdent	pdDiag	pdSymm
OR ₇	Var(β_1)	0,000441	23,8228	23,0968
	Var(β_2)	0,000441	197,6114	222,4675
	Var(β_3)	0,000441	0,000617	0,000654
	Corr(β_1, β_2)	–	–	-0,366
	Corr(β_1, β_3)	–	–	0,082
	Corr(β_2, β_3)	–	–	-0,477
	Var(ε)	30,7923	10,9068	10,8875
GO ₇	Var(β_1)	0,001659	16,9355	14,0642
	Var(β_2)	0,001659	128,4612	129,3036
	Var(β_3)	0,001659	0,001265	0,001013
	Corr(β_1, β_2)	–	–	0,007
	Corr(β_1, β_3)	–	–	0,954
	Corr(β_2, β_3)	–	–	-0,289
	Var(ε)	73,7549	46,6390	47,7650
LO ₇	Var(β_1)	0,000228	13,9589	14,5800
	Var(β_2)	0,000228	124,9638	135,2380
	Var(β_3)	0,000228	0,000145	0,000089
	Corr(β_1, β_2)	–	–	-0,222
	Corr(β_1, β_3)	–	–	0,988
	Corr(β_2, β_3)	–	–	-0,099
	Var(ε)	98,1762	68,0157	69,7637

Nota: Var(β_j) corresponde à variância dos efeitos aleatórios associados aos parâmetros β_1 , β_2 e β_3 . As correlações são apresentadas apenas para a estrutura geral positiva-definida (pdSymm). pdIdent = matriz identidade; pdDiag = matriz diagonal; pdSymm = matriz geral positiva-definida. OR – Ørskov e McDonald; GO – Gompertz; LO – Logístico.

ESTIMATIVAS DOS COMPONENTES DA MATRIZ DE COVARIÂNCIA RESIDUAL $\mathbf{R}$

Tabela 11 – Componentes estimados associados às diferentes estruturas da matriz de correlação residual ($\mathbf{R}$) nos modelos não lineares mistos.

Modelo	Estrutura	Parâmetro	$SD(\beta_1)$	$SD(\beta_2)$	$SD(\beta_3)$
OR _{7.2.6}	Independente	—	4,881	14,057	0,0248
	CS	$\rho \approx 0$	0,0008	10,979	0,0284
	CS + ID	$\rho = 0,568$	2,432	13,770	0,0232
	EXP	range = 47,40	0,0015	11,264	0,0192
	EXP + ID	range = 35,33	0,0649	1,736	0,0465
	GAUS	range = 3,69	0,0385	3,614	0,0465
	GAUS + ID	range = 3,52	0,0006	12,834	0,0247
GO _{7.2.6}	Independente	—	4,115	11,334	0,0356
	CS	$\rho \approx 0$	0,0027	0,0104	0,0500
	CS + ID	$\rho = -0,053$	4,964	11,256	0,0342
	EXP	range = 107,53	0,0053	11,524	0,0220
	EXP + ID	range = 102,04	0,0237	11,652	0,0000
	GAUS	range = 4,48	0,0649	1,736	0,0465
	GAUS + ID	range = 4,39	0,0001	11,201	0,0000
LO _{7.2.6}	Independente	—	3,736	11,179	0,0121
	CS	$\rho \approx 0$	0,0033	0,0118	0,0000
	CS + ID	$\rho = -0,053$	4,967	10,990	0,0116
	EXP	range = 99,42	0,0237	11,652	0,0000
	EXP + ID	range = 106,61	0,0112	12,938	0,0236
	GAUS	range = 4,82	0,0112	12,938	0,0236
	GAUS + ID	range = 4,79	0,0001	11,201	0,0000

Nota: $SD(\beta_j)$ corresponde ao desvio-padrão dos efeitos aleatórios associados aos parâmetros do modelo. CS = simetria composta; EXP = estrutura exponencial; GAUS= estrutura gaussiana. Modelos com “+ ID” indicam incorporação de heterogeneidade residual (função de variância). Valores de ρ próximos de zero indicam ausência de correlação residual relevante.

AJUSTES ADICIONAIS PARA FINS COMPARATIVOS

Tabela 12 – Indicadores de instabilidade estrutural no modelo OR₇ com **D** geral (pdSymm) e estruturas de variância GAUS + ID.

Componente	Estimativa
Correlação $\rho(\beta_1, \beta_3)$	-0,972
Correlação $\rho(\beta_2, \beta_3)$	-0,383
Correlação $\rho(\beta_1, \beta_2)$	0,165
Var(β_1)	0,9718
Var(β_2)	179,3929
Var(β_3)	0,000688
Variância residual	30,4308

Nota: Correlações próximas de ± 1 indicam forte dependência linear entre efeitos aleatórios, podendo sinalizar problemas de identificabilidade e instabilidade na estimação dos parâmetros.

RESULTADOS COMPLETOS DO TESTE DE TUKEY

Tabela 13 – Contrastes significativos no teste de Tukey ($p < 0,05$) para o parâmetro β_2 do modelo de Ørskov e McDonald (OR_{7.2.6}).

Contraste	Diferença	EP	gl	Teste t	Valor p (Tukey)
trat1 – trat9	-82,6755	12,4411	641	-6,6453	<0,0001
trat1 – trat15	-45,9076	12,4361	641	-3,6915	0,0217
trat2 – trat5	-39,4690	11,0684	641	-3,5659	0,0333
trat2 – trat6	-42,1982	11,0529	641	-3,8179	0,0139
trat2 – trat7	-40,9281	11,0380	641	-3,7079	0,0205
trat2 – trat8	-38,1526	11,0547	641	-3,4513	0,0482
trat2 – trat9	-90,7889	12,2229	641	-7,4278	<0,0001
trat2 – trat14	-52,4495	13,0458	641	-4,0204	0,0065
trat2 – trat15	-54,0210	12,2183	641	-4,4213	0,0012
trat3 – trat9	-62,8672	13,1875	641	-4,7672	0,0003
trat4 – trat9	-80,8709	12,1481	641	-6,6571	<0,0001
trat4 – trat15	-44,1030	12,1429	641	-3,6320	0,0267
trat5 – trat9	-51,3199	12,0116	641	-4,2725	0,0023
trat6 – trat9	-48,5907	11,9972	641	-4,0502	0,0058
trat7 – trat9	-49,8607	11,9836	641	-4,1608	0,0037
trat8 – trat9	-52,6363	11,9990	641	-4,3867	0,0014
trat9 – trat10	74,9175	13,1142	641	5,7127	<0,0001
trat9 – trat11	62,6300	12,1135	641	5,1703	<0,0001
trat9 – trat12	53,4907	11,9958	641	4,4591	0,0011
trat9 – trat13	82,1807	12,1208	641	6,7801	<0,0001
trat9 – trat16	58,4056	13,0508	641	4,4752	0,0010
trat13 – trat15	-45,4128	12,1156	641	-3,7483	0,0178