



Análise do Desempenho e Propagação do Erro em Modelos de Predição para Agricultura Digital Sustentável

Bianca O. Durgante^{1*}; Davi L. Lemos^{1†}; Matheus C. Correia^{1‡}
Naylor B. Perez^{2§}; Leonardo B. Pinho^{1¶}

Resumo: Em computação sustentável, a busca por soluções que reduzam o tempo de processamento está diretamente ligada à diminuição do consumo energético, fator essencial para o desenvolvimento de tecnologias mais eficientes. Este estudo investiga o desempenho computacional de modelos de Redes Neurais Recorrentes (RNN), com foco nas técnicas Long Short-Term Memory (LSTM) e Gated Recurrent Unit (GRU), aplicadas à Agricultura Digital (AD) no Pampa, com ênfase em Pecuária de Precisão. A partir de amostras coletadas entre 2014 e 2019, o desempenho dos modelos foi avaliado por diversas métricas, com ênfase no Erro Quadrático Médio (RMSE), comparando execuções com 1 e 12 núcleos disponíveis do processador, avaliando tanto a eficiência no tempo de execução quanto a acurácia das previsões. Os resultados indicam que a propagação do erro não seguiu um padrão linear, o que sugere a presença de ruído nos dados meteorológicos e a influência da aleatoriedade inerente aos modelos RNN. Além disso, observou-se que a eficiência computacional variou significativamente entre as técnicas e ambientes, destacando a importância de práticas de otimização tanto para o desempenho quanto para a sustentabilidade das tecnologias utilizadas em AD.

Palavras-chave: Inteligência Artificial; Pecuária de Precisão; Computação Verde; Desempenho Computacional; Propagação do Erro em Redes Neurais

*biancadurgante.aluno@unipampa.edu.br

†davilemos.aluno@unipampa.edu.br

‡matheuscorreia.aluno@unipampa.edu.br

§naylor.perez@embrapa.br

¶leonardopinho@unipampa.edu.br

¹UNIPAMPA – Campus Bagé – 96.413-172 – Bagé – RS – Brasil.

²Embrapa Pecuária Sul – 96.401-970 – Bagé – RS – Brasil.

1 INTRODUÇÃO

A integração entre tecnologia e agronegócio tem sido amplamente discutida na literatura, especialmente no contexto do aumento da produtividade e do melhor aproveitamento dos recursos naturais diante do agravamento das questões climáticas. Com o aumento populacional e a crescente demanda por alimentos, surgem grandes desafios como o de aumentar a produção agrícola sem ampliar significativamente a área plantada (Massruhá *et al.*, 2020) e nesse cenário técnicas de agricultura de precisão e práticas de agricultura digital surgem como importantes ferramentas de apoio para o produtor rural.

Em sistemas produtivos focados na pecuária extensiva de corte, a aplicação de técnicas de *smart farming* pode oferecer suporte essencial para o manejo das pastagens, otimizando o ganho de peso dos animais por meio de ajustes refinados nas taxas de lotação. Para otimizar a quantidade de animais em uma área de pastagem, é imprescindível estimar com precisão a quantidade de massa de forragem (MF), medida neste estudo pela taxa de acúmulo (TA) em quilogramas por hectare por dia (kg/ha/dia), disponível para o rebanho. É importante ressaltar que atualmente, nenhum método de predição oferece precisão absoluta, visto que diversos fatores ambientais, variações climáticas e imprecisões inerentes aos métodos de coleta de dados podem interferir na precisão, contudo as redes neurais devem ser capazes de prever a disponibilidade futura da massa de forragem com um nível satisfatório de acurácia. Essa capacidade de previsão é crucial para a gestão eficiente dos recursos, permitindo adotar o número de animais às condições disponíveis.

Essas previsões permitem direcionar os animais de maneira eficiente, evitando desperdício de recursos em períodos de abundância e prevenindo a perda de peso durante condições climáticas adversas, quando a produção de massa forrageira é menor e pode ser necessário recorrer à suplementação. Nesse sentido, foi desenvolvido por Schulte (2019) um modelo baseado em Redes Neurais Recorrentes (RNN) com arquitetura Long Short-Term Memory (LSTM), que viabiliza o aprendizado a partir de séries temporais (Le *et al.*, 2019), e que, com considerável acurácia, tem como saída a predição da disponibilidade de massa de forragem instantânea (MStotal) disponível em uma determinada área de pastagem. Posteriormente, o modelo foi melhorado por Soares (2021), quando, dentre outros aprimoramentos, passou a prever a taxa de acúmulo da massa de forragem (TA). Adicionalmente, a aplicação foi aprimorada com o acréscimo de técnicas de autoajuste de hiperparâmetros com fins de melhoria de acurácia, por Lemos *et al.* (2024). Apesar de demonstrarem significativa acurácia e enriquecimento neste processo de manejo, todos os modelos convergem a um mesmo déficit relativo ao seu re-treinamento, visto que o conjunto de dados utilizado para treinamento da aplicação deve ser constantemente atualizado e expandido, fazendo com que essa etapa seja executada várias vezes, elevando o custo computacional para manutenção desta tecnologia.

Com a crescente evolução dos ambientes computacionais e o agravamento das questões climáticas, é responsabilidade dos profissionais de computação, especialmente dos engenheiros, adotar práticas que minimizem o impacto ambiental de suas atividades e

incluir essa temática em suas ações. À medida que máquinas com maior capacidade computacional são utilizadas, de forma a fornecer o máximo de trabalho com o menor consumo energético (Merritt, 2024), ampliam-se as possibilidades de exploração voltadas à eficiência de tempo de execução em diferentes ambientes computacionais, oferecendo assim uma oportunidade concreta de reduzir a pegada de carbono da computação. Essa abordagem não apenas otimiza o tempo de execução, mas também desempenha um papel crucial na melhoria da eficiência energética e na promoção da sustentabilidade no setor tecnológico. Ademais, a integração de aspectos Ambientais, Sociais e de Governança (ESG) na área de tecnologia torna-se cada vez mais relevante, tendo em vista que, além de contribuir para a preservação ambiental, assegura que o desenvolvimento tecnológico esteja alinhado à urgência do combate às mudanças climáticas.

Nesse sentido, no presente trabalho analisa-se quantitativamente o crescimento do Erro Quadrático Médio (RMSE) da aplicação em função do erro acumulado pelo método das RNN, posto que, à medida que cada nova previsão é feita, ou seja, é gerada uma nova saída, existe uma margem de erro associada a essa predição, que, por sua vez, é medida através do RMSE. De acordo com Moura (2018), as camadas recorrentes recebem como entrada também a saída anterior e, nesse sentido, a hipótese é de que, à medida que mais previsões são realizadas, esses erros podem se acumular e se amplificar, resultando em previsões cada vez menos precisas. De maneira complementar, realiza-se uma comparação entre a arquitetura LSTM, tradicionalmente utilizada nos modelos iniciais, e a Gated Recurrent Unit (GRU), que tem se destacado na literatura por apresentar melhor desempenho em termos de tempo de execução e acurácia (Bao *et al.*, 2023). Para essa análise comparativa, além da avaliação do tempo de execução, os modelos foram submetidos a diferentes ambientes computacionais, com o objetivo de identificar qual técnica oferece maior eficiência considerando os princípios de sustentabilidade. Essa avaliação não apenas permite a adoção de práticas de computação verde, mas também contribui para a mitigação dos impactos ambientais associados à computação. Ao considerar o custo computacional e o tempo de execução das aplicações de Inteligência Artificial, promove-se uma maior conscientização sobre as implicações ecológicas das diferentes alternativas, potencializando os benefícios da computação sustentável.

O restante deste artigo é segmentado em cinco seções: a Seção 2 é relativa aos trabalhos correlatos, explanando brevemente o histórico de trabalhos com o mesmo eixo temático e a revisão de literatura realizada; a Seção 3 explora a questão dos diferentes ambientes computacionais de execução; a Seção 4 diz respeito à metodologia, apresentando materiais e métodos aplicados para testes e análises; a Seção 5 reúne os resultados e discussões, sintetizando os dados levantados e comparando-os e, por fim, a Seção 6 sintetiza as conclusões e destaca trabalhos futuros.

2 TRABALHOS RELACIONADOS

Existem diversas técnicas para estimar a massa de forragem de uma área de pastagem. Dentre elas, destaca-se o método direto que pode ser realizado através do corte utilizando gaiolas para isolamento de uma área específica onde acontece o acúmulo da forragem (Gardner, 1986). Tal abordagem foi utilizada para preencher os dados reais que compõem as amostras do modelo que é o objeto de estudo deste trabalho.

Schulte (2019) apresenta um modelo de predição de disponibilidade da oferta de massa de forragem (M_{total}) em pastagem natural denominado Modelo Original, baseado em RNN de arquitetura LSTM, neste caso, com uma estrutura estabelecida de 30 e 15 neurônios para a primeira e a segunda camada, respectivamente. Esta rede é composta por dados amostrais reais do bioma Pampa coletados por meio do método direto, e dados meteorológicos. Esta aplicação foi capaz de estimar com significativa acurácia, quando comparado às previsões das abordagens adotadas anteriormente, contribuindo assim com técnicas de *smart farming* como ferramentas de suporte à decisão em pastagens.

Adicionalmente, Soares (2021) aprimorou o Modelo Original ao incorporar dados de previsão meteorológica para a predição, mantendo o ajuste empírico dos hiperparâmetros realizado no primeiro modelo e projetando o TouceiraTech, um protótipo de *Farm Management Information Systems* para Pecuária de Precisão. Além disso, em sua pesquisa, a predição da oferta de M_{total} foi substituída pela predição do acúmulo de forragem em relação ao mês anterior, caracterizando TA. Houve também uma alteração na captura do efeito estacional, o qual passou a ser representado pelo mês e ano, ao invés da data específica da amostra, a nova aplicação denomina-se Modelo Ajustado.

Com o objetivo de aprimorar a acurácia da predição do Modelo Original, Lemos *et al.* (2024) aplicou o KerasTuner (O'Malley *et al.*, 2019), ferramenta de autoajuste de hiperparâmetros, utilizando-se do ambiente de execução Google Colaboratory (Colab) em sua versão gratuita, que contava com Unidade Central de Processamento (CPU), Unidade de Processamento Gráfico (GPU) e Unidade de Processamento de Tensor (TPU), e possui nos respectivos ambientes, processador Intel Xeon 2,2 GHz utilizando de um núcleo com duas threads (possuindo, em sua totalidade, dez núcleos com vinte threads), GPU Tesla T4 15 GB GDDR6 com 2560 núcleos e uma TPU v2 16 GB High Bandwidth Memory (HBM) com oito núcleos, denominando-se o modelo resultante da aplicação da ferramenta de autoajuste, Modelo Autoajustado. A partir dos resultados obtidos foram incorporados aspectos de ESG ao quantificar os impactos da aplicação do método em relação ao tempo de execução, considerando a etapa adicional de sintonização do uso da ferramenta, conforme ilustra a Tabela 1, que traz o tempo de execução dos modelos em diferentes ambientes computacionais do Colab.

Tabela 1: Tempo de execução em diferentes ambientes do Colab

Ambiente	Tempo de execução (s)	
	Modelo Original	Modelo Original Autoajustado
CPU	267,62	786,02
GPU	214,71	601,14
TPU	275,49	841,32

Apesar disso, na revisão de literatura realizada, não foram encontrados trabalhos correlatos que abordem especificamente a análise do acúmulo de erro decorrente do uso de algoritmos de RNN com a comparação de desempenho em diferentes ambientes computacionais no contexto da pecuária. Esse achado destaca a relevância e a contribuição potencial desse estudo para o avanço do conhecimento na área, preenchendo uma lacuna significativa existente na pesquisa.

3 EXPLORAÇÃO DE AMBIENTES COMPUTACIONAIS

Com o objetivo de realizar um estudo mais específico sobre o tempo de execução dos modelos em diferentes ambientes computacionais, este artigo, em continuidade ao trabalho desenvolvido por Lemos *et al.* (2024), dedica-se a quantificar o desempenho dos modelos quando submetidos a condições variáveis de recursos computacionais. Da mesma forma, dando continuidade a esta linha de pesquisa, explora-se o aspecto sugerido anteriormente envolvendo a utilização da arquitetura GRU. Além disso, a análise reforça os aspectos relacionados à ESG, visando promover uma abordagem sustentável no uso de recursos computacionais de forma a identificar a abordagem mais eficiente para o processo de re-treinamento do modelo.

Em geral, técnicas de aprendizado profundo, como as aplicadas nesta pesquisa, requerem a disponibilização de recursos computacionais para uso e possibilidade de otimização no tempo de execução da aplicação. No que tange o sistema aplicado nos experimentos deste estudo, destaca-se, em execução local, o uso do microprocessador multicore AMD Ryzen 5 4500 com frequência de 3,6 GHz e seis núcleos com doze *threads*, placa de vídeo AMD Radeon RX 6600 8 GB GDDR6 com 1792 CUDA *cores*, 32 GB DDR4 com 3200 MHz e Sistema Operacional Ubuntu 24.04 LTS.

Para uma melhor distinção das etapas contidas dentro de uma execução dos modelos, considerando a relevância dessa informação para a interpretação dos resultados trazidos na Seção 5, destaca-se que estas são divididas e analisadas em seis partes principais, abstraindo etapas de sintetização: manipulação dos dados, definição e treinamento do modelo original, cálculo de métricas do modelo original, definição e busca aleatória do KerasTuner, treinamento do modelo com KerasTuner e cálculo das métricas.

Com o objetivo de avaliar o impacto do *hardware* na execução dos algoritmos, foram utilizados um e doze núcleos virtuais do processador local mencionado. O experimento foi repetido para as arquiteturas LSTM e GRU, e o desempenho foi medido pelo tempo de

execução, calculado a partir dos *logs* de saída das aplicações, que registram o tempo gasto em cada etapa do modelo. A soma dessas etapas resultou no tempo total de execução. Os resultados foram trazidos juntamente com aqueles obtidos por meio da execução no ambiente Colab, baseados na experimentação realizada por Lemos *et al.* (2024). É importante ressaltar que não se pretende utilizar esse ambiente para a execução do modelo devido à necessidade periódica de retreinamento. No entanto, sua contribuição se dá na fase de aprimoramento das aplicações, em virtude das vantagens associadas ao uso de ambientes de execução compartilhada. Destaca-se que neste trabalho, o autoajuste foi aplicado ao Modelo Adaptado, aprimorado por Soares (2021), que resultou no acréscimo de atributos de entrada e aumento do número de amostras em relação ao Modelo Original. Como resultado, espera-se teoricamente um maior tempo de treinamento quando os modelos Original e Adaptado são comparados em um mesmo ambiente computacional.

4 METODOLOGIA

Esta pesquisa exploratória experimental consiste no aprimoramento de um modelo para suporte à decisão para manejo de pastagens originalmente definido por Schulte (2019) e aprimorado por Soares (2021). Também é utilizado o modelo apresentado por Lemos *et al.* (2024), que aplica a ferramenta de autoajuste de hiperparâmetros KerasTuner, sintonizada com o parâmetro *RandomSearch* (busca aleatória), onde, pela primeira vez, abriu-se mão do ajuste empírico aplicado nos modelos anteriores, visto que já demonstrava grande potencial para aprimoramento na acurácia das predições realizadas. Neste estudo, além dos modelos Adaptado e Autoajustado desenvolvidos nos trabalhos citados anteriormente, utiliza-se também a aplicação da técnica GRU com os mesmos parâmetros de configuração dos demais modelos.

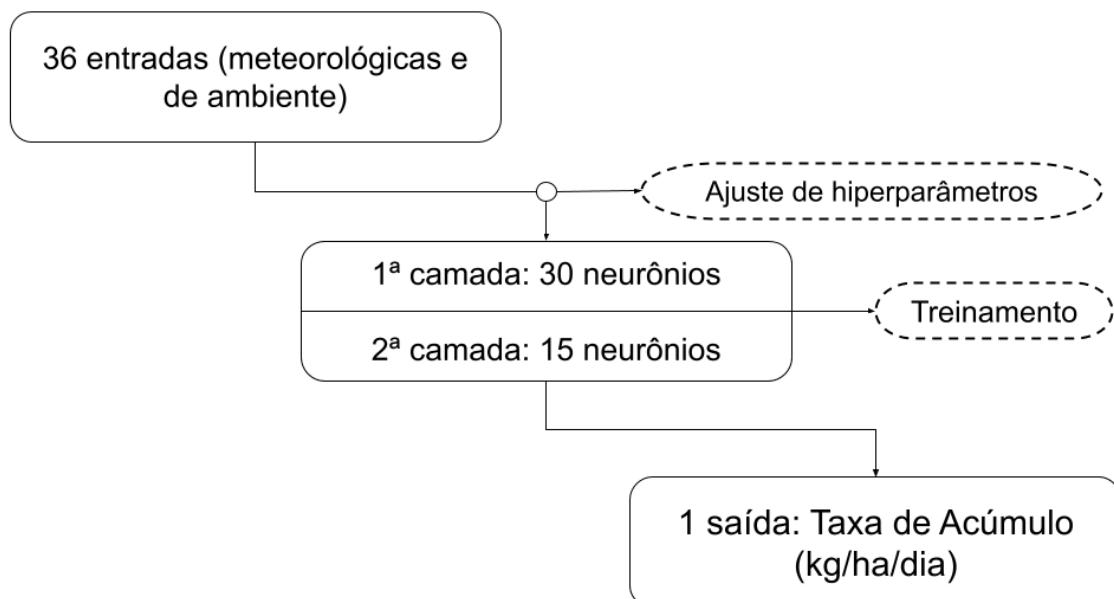


Figura 1: Fluxo de funcionamento do modelo

O funcionamento dos modelos segue três etapas principais, conforme ilustrado na Figura 1, utilizando RNN para processar séries temporais. As principais bibliotecas Python que compõem as aplicações são Keras, Pandas, Tensorflow, Scikit-learn e NumPy. O conjunto de dados em questão reúne 48 amostras coletadas entre janeiro de 2014 e agosto de 2019, sendo que cada uma é composta por 36 atributos de entrada que consistem em séries temporais estratificadas sobre uma delimitação geográfica e em dados meteorológicos, extraídas de um banco de dados espacial, bem como uma saída referente à TA. A técnica aplicada fornece a quantidade de MF isolada do consumo e influência animal para uma determinada área (DG), possibilitando contraste com a quantidade observada em uma área equivalente submetida aos impactos do pastejo (FG), de forma que estabeleça-se a relação ilustrada pela equação (1), onde ΔT é a variação de dias entre as medições.

$$TA = \frac{DG - FG}{\Delta T} \quad (1)$$

A fim de viabilizar a mensuração de se há impacto relacionado ao acúmulo de erro, ou não, na predição de vários meses pelos modelos, o caso definido para estudo é o do último mês coletado (Agosto/2019), garantindo que este seja predito em todos os testes realizados. Desta forma, utilizou-se doze diferentes alocações para o conjunto de teste, de forma que variasse da predição mensal até a predição anual, sendo o caso de estudo sempre o último a ser predito pela RNN. Reconhecendo-se a aleatoriedade das RNN, fator ainda mais presente no caso do KerasTuner quando configurado para sintonização através de busca aleatória, repetiu-se para cada variação de tamanho do conjunto de testes, dez vezes cada modelo, na tentativa de mitigar tal randomicidade.

Para avaliar o desempenho dos modelos perante os resultados dos treinamentos e testes, foram escolhidas três métricas de avaliação principais: RMSE, desvio padrão e variância. O RMSE é uma métrica estatística utilizada para avaliar a acurácia de modelos preditivos e calculado a partir da média das magnitudes dos erros de predição, onde o erro quadrático médio é definido como a diferença entre o valor observado e o valor predito pelo modelo. O RMSE é calculado conforme a equação (2).

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (2)$$

Onde y_i é o valor real, $\hat{y}_i$ o valor predito e n a quantidade de predições.

O desvio padrão (σ) dos erros de predição indica quão dispersos estão os erros em relação à média dos erros, nesse sentido, um σ menor sugere que o modelo tem uma maior consistência nos resultados preditos. O σ é calculado conforme a equação (3).

$$\sigma = \sqrt{\frac{1}{n} \sum_{i=1}^n (e_i - \bar{e})^2} \quad (3)$$

Onde $e_i = y_i - \hat{y}_i$ é o erro de predição, ou seja, a diferença entre o valor real y_i e o valor predito $\hat{y}_i$. A média dos erros de predição é $\bar{e}$ e n é a quantidade de meses preditos.

A variância é o quadrado do desvio padrão e oferece uma medida de quão dispersos estão os erros de predição. Neste sentido, a variância avalia a variação dos erros em torno da média dos erros das predições do modelo, sendo caracterizada pela equação (4).

$$\sigma^2 = \frac{1}{n} \sum_{i=1}^n (e_i - \bar{e})^2 \quad (4)$$

5 RESULTADOS E DISCUSSÕES

Tendo em perspectiva a avaliação do impacto dos diferentes ambientes computacionais na execução destes modelos, a execução dos testes descritos na Seção 3 têm, através das médias dos tempos de execução (em segundos) de todas as diferentes configurações e ambientes de execução, seus resultados expressados na Figura 2. Tendo em vista a análise do impacto ocasionado pela quantidade de núcleos disponíveis e a exploração do paralelismo pela RNN, observa-se aproveitamento considerável ao comparar-se as execuções com um e com doze núcleos do processador local, trazendo, no caso da execução com doze núcleos, melhoria de aproximadamente 35% para o Modelo Ajustado e de 48% para o Modelo Autoajustado (ambos com arquitetura LSTM). Destaca-se que a denominação Autoajustado dos modelos refere-se ao ajuste dos hiperparâmetros realizados pela ferramenta KerasTuner. Bem como os modelos Original e Ajustado sendo as aplicações e melhoramentos mencionadas na Seção 2.

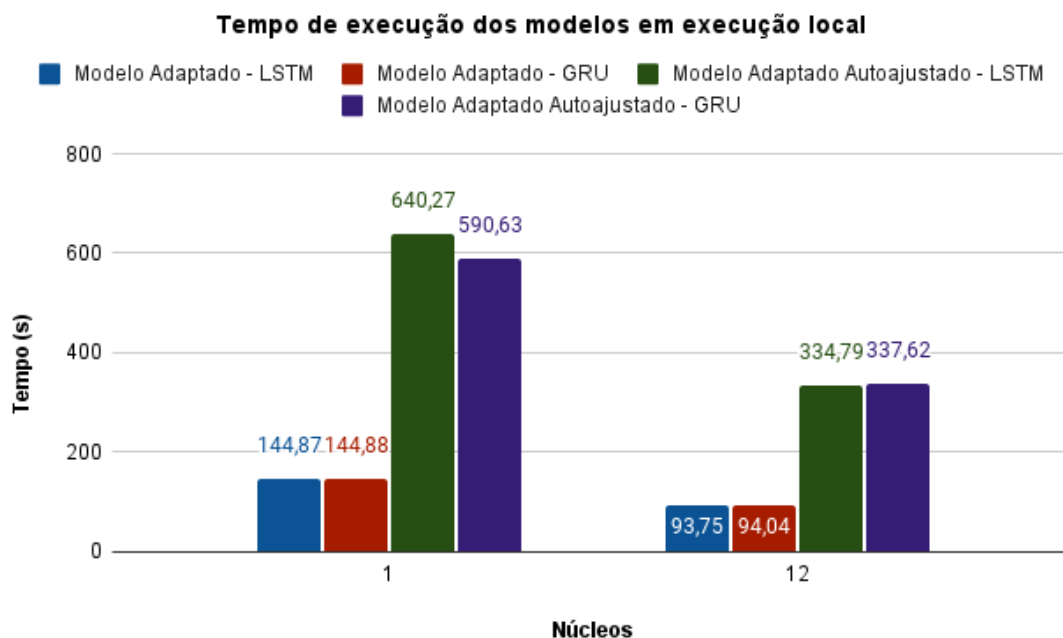


Figura 2: Resultados dos modelos em diferentes ambientes computacionais

Comparando-se as duas técnicas de RNN utilizadas, a análise do tempo de execução em diferentes combinações dos modelos e ambientes de execução não evidenciou a superioridade consistente da GRU em termos de velocidade, segundo esperado com base na literatura. Conforme a Figura 2, os tempos de execução dos modelos foram similares na maioria dos cenários, contrariando a expectativa de que a GRU apresentaria vantagens significativas em relação à LSTM. No entanto, destaca-se o desempenho da GRU em um ambiente com o processador restrito a um único núcleo, especificamente quando autoajustado, onde houve uma redução de tempo de aproximadamente 10% em comparação à LSTM sob as mesmas condições. Isso sugere que, em ambientes com recursos computacionais limitados, a GRU pode explorar de maneira mais eficiente os recursos disponíveis, apresentando um desempenho superior à LSTM nessas circunstâncias.

A Figura 3 explora o comportamento do RMSE ao aplicar-se as duas técnicas de RNN (LSTM e GRU), no Modelo Autoajustado. Com base nos resultados obtidos, especula-se uma afinidade entre a ferramenta KerasTuner e a arquitetura GRU. Ao analisar os valores de RMSE gerados ao longo de dez execuções do modelo de predição, observa-se que a arquitetura LSTM apresenta maior oscilação em relação ao RMSE médio, com um σ de 10,54. Em contrapartida, GRU demonstra um comportamento mais estável, com menor variação em torno da média, evidenciado por um σ de 4,27. Esses resultados sugerem uma maior consistência da GRU na minimização do RMSE, possivelmente indicando uma melhor afinidade com o processo de autoajuste proporcionado pelo KerasTuner.

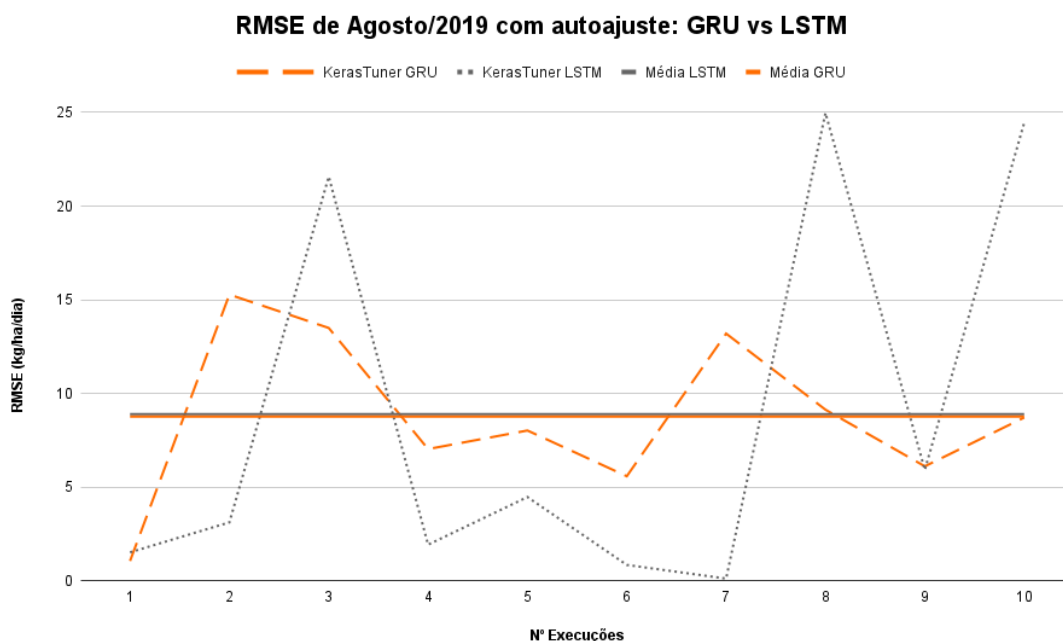


Figura 3: Gráfico do RMSE para os modelos autoajustados com arquitetura GRU e LSTM

Quanto à análise do acúmulo de erro, tendo como referência o último mês da série histórica (Agosto/2019), os resultados observados nos testes encontram-se agrupados na Tabela 2, a qual aponta que o Modelo Adaptado (ajustado empiricamente) apresenta

desempenho superior ao Modelo Autoajustado (ajustado com KerasTuner) em termos de RMSE, desvio padrão e variância. O RMSE médio do Modelo Adaptado é mais baixo, 4,73 kg/ha/dia, em contraste com 10,18 kg/ha/dia do Modelo Autoajustado. O desvio padrão também é menor no Modelo Adaptado, com 2,66, contra 2,92 do Modelo Autoajustado, evidenciando a superioridade. A variância do ajustado empiricamente é 6,48, comparada a 7,83 do ajustado com KerasTuner, indicando previsões menos dispersas em relação aos valores reais. Esses resultados indicam que o ajuste empírico se mostra mais eficaz, oferecendo maior acurácia nas previsões da taxa de acúmulo.

Resultados preliminares levantaram a hipótese de que KerasTuner tenha seu desempenho diretamente impactado quando há alteração do formato de captura do efeito sazonal do modelo. Dentre todas as combinações possíveis, o uso da ferramenta de autoajuste acaba apenas se sobressaindo aos resultados obtidos pelo ajuste empírico quando a variável relativa à data está no molde "dia-mês-ano" (Modelo Original), diferentemente do desempenho obtido com "mês" e "ano" (sendo variáveis distintas para entrada), como é o caso do Modelo Adaptado onde o uso de autoajuste não representa melhoria para acurácia. Portanto, destaca-se que esta particularidade pode afetar também a análise de acúmulo de erro.

Tabela 2: Resultados constatados para o mês de Agosto/2019

Quantidade de Meses Preditos	Taxa de acúmulo (kg/ha/dia)		RMSE		
	Real	Predita (média)		Modelo Adaptado	Modelo Autoajustado
		Modelo Adaptado	Modelo Autoajustado		
1	29,77	35,55	38,83	5,78	9,05
2		38,79	18,24	9,02	11,52
3		30,78	17,72	1,01	12,05
4		36,23	19,32	6,46	10,45
5		33,77	16,35	4,00	13,42
6		31,80	14,84	2,03	14,93
7		33,47	17,44	3,70	12,33
8		34,15	18,86	4,38	10,91
9		32,59	21,98	2,82	7,79
10		31,86	22,98	2,09	6,79
11		36,18	24,76	6,41	5,01
12		38,87	21,86	9,10	7,91
MÉDIA				4,73	10,18
VARIÂNCIA				6,48	7,83
DESVIO PADRÃO				2,66	2,92

Adicionalmente, os dados são apresentados graficamente na Figura 4, onde são correlacionadas a quantidade de previsões realizadas com a variação do RMSE da previsão para o mês de agosto de 2019, visando proporcionar uma análise visual do

comportamento do RMSE à medida que aumenta o número de meses preditos. Dessa forma, destacam-se os meses com os melhores e piores desempenhos nas previsões da RNN. No contexto do Modelo Adaptado, que incorpora o ajuste empírico dos hiperparâmetros, observa-se que o melhor desempenho ocorre ao prever três meses (RMSE de 1,01 kg/ha/dia), enquanto o pior desempenho é registrado ao prever doze meses (RMSE de 9,1 kg/ha/dia). Para o Modelo Autoajustado, os melhores e piores desempenhos são observados ao prever onze meses (RMSE de 5,01 kg/ha/dia) e seis meses (RMSE de 14,93 kg/ha/dia), respectivamente.

A partir da Figura 4, que analisa isoladamente o comportamento do RMSE para o mês de agosto de 2019 ao longo das previsões de um ano completo, buscou-se identificar a existência de um aumento do erro conforme o número de previsões cresce. Os resultados indicam que não há um crescimento linear no valor do RMSE, conforme poderia ser esperado, dado que as novas previsões são utilizadas como entradas, em conformidade com a definição do modelo (Moura, 2018).

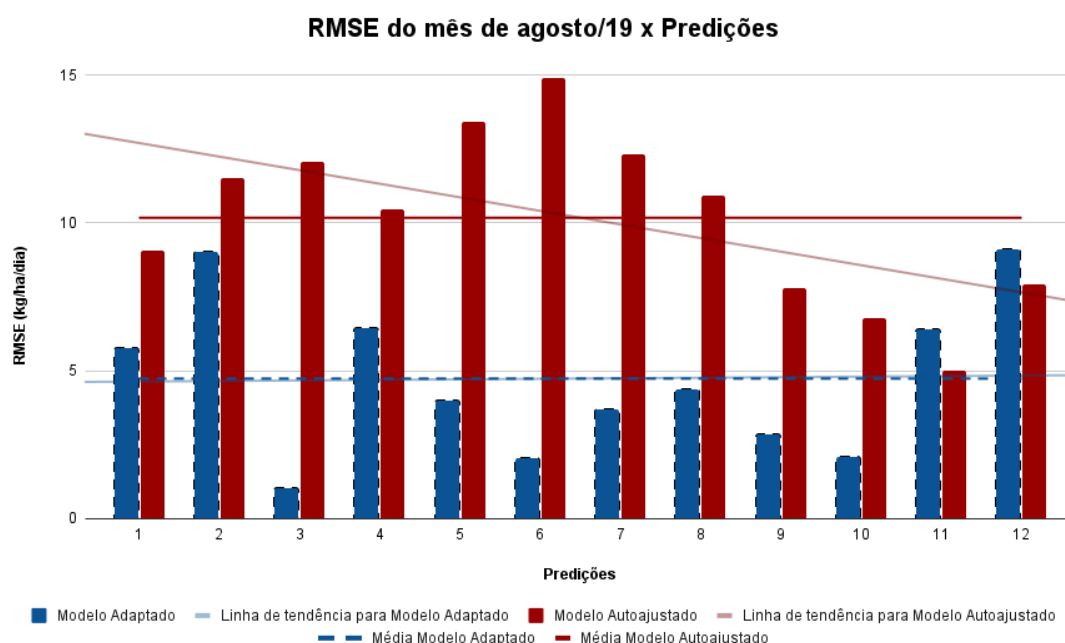


Figura 4: Relação entre a quantidade de meses preditos e o erro da previsão do último mês

Na avaliação do Modelo Autoajustado, observa-se que o RMSE mais alto (pior caso) ocorre ao prever seis meses, ao contrário da hipótese inicial, que indicava um desempenho inferior nas previsões de doze meses. Diante desse resultado, foram realizadas análises adicionais para investigar os fatores que podem ter influenciado a ferramenta de autoajuste, mesmo que não diretamente, relacionados ao aumento do número de previsões.

Ao analisar-se o RMSE do mês de agosto de 2019 em relação à sua variação ao longo das doze meses de previsão, torna-se evidente que não há um padrão definido no comportamento desse aumento. Diante dessa constatação, foram avaliadas as entradas do modelo, observando-se que muitas dessas variáveis são meteorológicas. Nesse contexto,

levanta-se a hipótese de que a imprevisibilidade climática, acentuada por mudanças cada vez mais severas, pode ter interferido negativamente no aprendizado e nas predições do modelo. Isso reforça a necessidade do desenvolvimento de modelos mais resilientes às mudanças climáticas, de forma que estes sejam capazes de considerar eventos atípicos de maneira que impactem o processo de aprendizagem de forma mais eficaz.

6 CONCLUSÃO

Este estudo investigou o acúmulo do erro em modelos de predição utilizando diferentes técnicas de RNN, aplicados na previsão da taxa de acúmulo de massa de forragem no bioma Pampa, adicionalmente avaliando o desempenho destes em diferentes ambientes computacionais. Os resultados destacaram diferenças significativas entre o Modelo Ajustado empiricamente e o Modelo Autoajustado por KerasTuner, ressaltando ainda, de forma gráfica, a variação do RMSE do mês de agosto de 2019 ao longo das predições de um ano inteiro. Ao explorar as arquiteturas LSTM e GRU os resultados mostraram o potencial para uso de GRU quando combinada ao uso da ferramenta KerasTuner com destaque para melhor aproveitamento de um *hardware* inferior.

Para o modelo ajustado empiricamente, observou-se que o menor RMSE ocorreu quando foram previstos três meses, enquanto o maior erro foi encontrado na previsão de doze meses. Esses resultados indicam que, embora o modelo ajustado empiricamente possa fornecer previsões precisas em curto prazo, sua precisão diminui significativamente com o aumento do período de previsão. Por outro lado, o Modelo Autoajustado por KerasTuner apresentou seu melhor desempenho na previsão de onze meses e o pior desempenho na de seis meses. Esse comportamento sugere que o Modelo Autoajustado pode lidar melhor com previsões de longo prazo, embora mostre variações mais significativas no curto prazo devido à aleatoriedade da busca por hiperparâmetros e da abordagem das RNN.

No contexto da análise de diferentes ambientes computacionais e arquiteturas de RNN, considera-se a hipótese de que a GRU é mais eficaz na exploração de *hardwares* menos potentes e, através disso, maximiza o uso de recursos computacionais e emerge como uma alternativa promissora para desenvolver modelos mais eficientes energeticamente, auxiliando na redução dos custos de re-treinamento. O uso dessa arquitetura combinada a técnicas de virtualização, pode permitir a execução de diversos modelos de predição dentro de um mesmo ambiente computacional, graças ao seu baixo custo de processamento. Essa abordagem, quando combinada a outras ferramentas de suporte à decisão e manejo, potencializa a aplicação de técnicas de agricultura digital no dia-a-dia do produtor.

De forma geral, pode-se concluir, com base nos dados e resultados dos experimentos, que não há um impacto significativo causado pelo acúmulo de erro nas predições realizadas pelos algoritmos de RNN. Essa conclusão contraria a hipótese inicial, onde esperava-se que o último mês a ser predito teria a pior acurácia e seria possível identificar um comportamento linear no crescimento do RMSE. Da mesma forma, observou-se que a GRU não resultou em uma redução geral no tempo de execução do modelo de predição. No

entanto, demonstrou um potencial significativo na diminuição da demanda computacional, mesmo quando utilizada em conjunto com a ferramenta de autoajuste de hiperparâmetros. Essa característica ressalta a contribuição da GRU para práticas de computação verde, ao permitir um uso mais eficiente dos recursos disponíveis, alinhando-se com os objetivos de sustentabilidade e eficiência energética na execução de modelos preditivos.

Em suma, esta pesquisa visa contribuir para o estado da arte, incorporando aspectos Ambientais, Sociais e de Governança (ESG), análises estatísticas, *Smart Farming* e tecnologias modernas na agricultura digital e pecuária de precisão. Além disso, destaca-se a intenção de pesquisas e trabalhos futuros voltados ao aprimoramento da técnica e das variáveis do modelo. Entre essas iniciativas, pretende-se incorporar a previsão de cenários climáticos extremos e analisar seu impacto no modelo, visando aumentar a robustez do algoritmo frente a esses fatores. Isso permitirá que dados aberrantes, resultantes da maior frequência desses eventos, sejam considerados, embora com um peso diferente em relação às demais variáveis de entrada, garantindo a coerência dos resultados de predição. Também busca-se aprimorar o baixo custo computacional da arquitetura GRU por meio da adoção do método de captura estacionária do Modelo Original, que melhora a compatibilidade do KerasTuner com a acurácia. Complementarmente, destaca-se a intenção de explorar técnicas de virtualização local com base nos resultados obtidos, nos quais uma das combinações testadas demonstrou potencial para oferecer uma boa relação entre acurácia e tempo de execução utilizando menos recursos computacionais. Isso permitirá a adoção de práticas de computação sustentável ao potencializar o aproveitamento dos recursos da máquina local, possibilitando a execução simultânea de múltiplas aplicações.

Outro aspecto relevante é a exploração das temáticas de computação verde, analisando a capacidade de aproveitamento de técnicas de paralelismo em modelos RNN para melhorar o tempo de execução, reduzir o consumo de energia e aumentar a eficiência computacional, alinhando o desenvolvimento tecnológico com princípios de sustentabilidade ambiental. Além disso, ressalta-se que já estão sendo desenvolvidos trabalhos que buscam reconhecer padrões comportamentais de sintonização da rede, estabelecendo uma relação direta entre hiperparâmetros, acurácia e erro dos algoritmos. Posto que esse ajuste empírico mais adequado, com um viés de ESG, pode mitigar o processo custoso de sintonização e o gasto de tempo em treinamentos com resultados insatisfatórios. Por fim, métodos de instrumentação dos modelos analisados também estão sendo implementados de forma a detalhar a demanda de recursos computacionais da aplicação.

AGRADECIMENTOS

Reconhece-se o apoio da CAPES - Código de Financiamento 001, como parte das ações integradas do Programa de Pós-Graduação em Computação Aplicada (PPGCAP) com Ensino, Pesquisa e Extensão de Engenharia de Computação na graduação, além dos Programas Institucionais de Bolsas de Iniciação Científica PIBIC/CNPq, PROBIC/FAPERGS e PRO-IC/UNIPAMPA e do Programa de Demanda Social da CAPES.

REFERÊNCIAS

- Bao, K., J. Bi, R. Ma, Y. Sun, W. Zhang, e Y. Wang (2023). “A Spatial-Reduction Attention-Based BiGRU Network for Water Level Prediction”. Em: *Water* 15, p. 1306. DOI: 10.3390/w15071306.
- Gardner, A. (1986). *Técnicas de pesquisa em pastagens e aplicabilidade de resultados em sistemas de produção*. Vol. 1. INCA/EMBRAPA, p. 197.
- Le, X.-H., H. V. Ho, G. Lee, e S. Jung (2019). “Application of Long Short-Term Memory (LSTM) Neural Network for Flood Forecasting”. Em: *Water* 11.7. DOI: 10.3390/w11071387. URL: <https://www.mdpi.com/2073-4441/11/7/1387>.
- Lemos, D., B. Durgante, N. Perez, e L. Pinho (2024). “Impactos da Ferramenta KerasTuner em Modelo de Predição Baseado em LSTM para Pecuária Sustentável”. Em: *Anais da XXIV Escola Regional de Alto Desempenho da Região Sul*. Florianópolis/SC: SBC, pp. 97–100. DOI: 10.5753/erads.2024.238568. URL: <https://sol.sbc.org.br/index.php/erads/article/view/28012>.
- Massruhá, S. M. F. S., M. A. de A. Leite, S. R. de M. Oliveira, C. A. A. Meira, A. L. Junior, e E. L. Bolfe (2020). *Agricultura Digital: Pesquisa, Desenvolvimento e Inovação nas Cadeias Produtivas*. Embrapa, pp. 20–45. URL: <https://www.embrapa.br/busca-de-publicacoes/-/publicacao/1126213/agricultura-digital-pesquisa-desenvolvimento-e-inovacao-nas-cadeias-produtivas>.
- Merritt, R. (2024). “O Que É Computação Verde?” Em: *Blog oficial NVIDIA Brasil*.
- Moura, G. B. (2018). “Redes neurais recorrentes para a classificação de estruturas retóricas”. Diss. de mest. Mestrado em Ciência da Computação - Universidade Estadual de Maringá.
- O'Malley, T., E. Bursztein, J. Long, F. Chollet, H. Jin, L. Invernizzi, *et al.* (2019). *KerasTuner*. <https://github.com/keras-team/keras-tuner>. [online: acesso em 24-janeiro-2024].
- Schulte, L. G. (2019). “Suporte à Decisão em Pastagens: Análise Espaço-temporal e Aprendizado de Máquina para Predição da Disponibilidade de Forragem no Contexto de Smart Farming”. Diss. de mest. Universidade Federal do Pampa (UNIPAMPA) – Programa de Pós-Graduação em Computação Aplicada (PPGCAP).
- Soares, A. F. (2021). “TouceiraTech: um Farm Management Information System para pecuária de precisão baseado em predição com redes neurais recorrentes”. Diss. de mest. Universidade Federal do Pampa (UNIPAMPA) – Programa de Pós-Graduação em Computação Aplicada (PPGCAP).