

ARTIGO ORIGINAL

Aplicação de estratégias de aprendizado de máquina na detecção de plantas invasoras em pastagens

Application of machine learning strategies in the detection of invasive plants in pastures

Gabriel Tadioto Oliveira ¹ and André Luiz Brun ¹

¹Universidade Estadual do Oeste do Paraná - UNIOESTE

tadiotogabriel@hotmail.com, andre.brun@unioeste.br

Recebido: 20/09/2024. Revisado: 12/07/2025. Aceito: 20/07/2025.

Resumo

A presença de plantas invasoras em pastagens é um desafio enfrentado pelo pecuarista. A erva daninha, além de concorrer com a forrageira por nutrientes, pode causar mal aos animais. O combate químico é uma das opções indicadas para enfrentar esse problema. Nesse sentido, a detecção automática é uma alternativa necessária, ao aumentar a agilidade e a diminuição dos custos do processo. Nesse sentido, algoritmos de aprendizado de máquina mostram-se como possíveis alternativas. Neste trabalho foram avaliados diferentes modelos para a classificação de plantas invasoras, tendo como referência imagens do gênero *Rumex*. Os modelos avaliados eram baseados em estratégias monolíticas e em sistemas de múltiplos classificadores. Os experimentos, realizados sobre dois conjuntos de dados com diferentes proporções entre as classes positivas e negativas, mostraram que o KNN foi o modelo individual mais competente em detectar a planta invasora, obtendo acurácia superior a 95% nos dois conjuntos e sensibilidade superior a 0,90. Já entre os modelos baseados em *ensembles* aquele que se destacou foi o *Random Forest*, com taxa de acertos acima de 0,95 e sensibilidade de aproximadamente 90% nos dois cenários.

Palavras-Chave: Classificação Automática; Ervas daninhas; Forrageiras; Machine Learning; Pecuária.

Abstract

The presence of invasive plants on pastures is a challenge for livestock farmers. Weeds not only compete with forage for nutrients, but can also harm livestock. Chemical control is one of the options to tackle this problem. In this sense, automatic detection is a necessary alternative, as it increases agility and reduces process costs. In this sense, machine learning algorithms emerge as possible alternatives. In this work, different models for the classification of invasive plants were evaluated, using images of the genus *Rumex* as reference. The evaluated models were based on monolithic strategies and on systems with multiple classifiers. The experiments, performed on two datasets with different proportions of positive and negative classes, showed that KNN was the most competent single model in detecting the invasive plant, achieving an accuracy of more than 95% and a sensitivity of more than 0.90 in both datasets. Among the models based on *ensembles*, *Random Forest* stood out with a accuracy of over 0.95 and a sensitivity of about 90% in both scenarios.

Keywords: Automatic Classification; Weeds; Forage; Machine Learning; Livestock.

1 Introdução

É notório que o Brasil tem um importante papel na pecuária mundial. O rebanho bovino brasileiro, segundo dados da Organização das Nações Unidas para a Alimentação e a Agricultura, foi o maior do mundo em 2020, com cerca de 217 milhões de cabeças (Aragão and Contini, 2022) e, segundo dados do IBGE, o tamanho dos rebanhos chegou a mais de 234 milhões de cabeças em 2022 (IBGE, 2022). Além de que, a área destinada à pecuária (pastagens) foi de 167 milhões de hectares onde se produziu cerca de 95% da carne bovina nacional (Portal do Agronegócio, 2022).

Um dos grandes desafios do produtor pecuarista é a presença de plantas invasoras. Ervas daninhas em meio às pastagens, competem diretamente por espaço, água, nutrientes e luz, podendo acarretar diversos danos, além de outros problemas indiretos, como envenenamento e ferimento dos animais, redução na qualidade do leite e aumento do tempo para a formação de pastagens (Pereira et al., 2011; Vilar, 2021).

Em um estudo de 2013, o professor Sidnei Roberto de Marchi identificou que a presença de plantas urticantes, ou com espinhos, fazem com que o gado evite pastejar a até um metro de distância, evidenciando a perda de espaço devido as ervas daninhas. Em relação à água, é comum verificar, em dias quentes, as plantas forrageiras murchando enquanto as invasoras não apresentam sinais de déficit hídrico. Um dos fatores é a taxa de exploração de volume do solo pelo sistema radicular, como ocorre com a ciganinha e o assa peixe (Victoria Filho et al., 2014).

Acerca da competição por nutrientes, o nitrogênio, o fósforo e o potássio são os mais importantes para o processo de competição. A planta de mostarda-brava (*Brassica campestris*), por exemplo, necessita duas vezes mais nitrogênio e fósforo e quatro vezes mais potássio que uma planta de aveia cultivada (Pereira et al., 2011), aumentando ainda mais a competitividade e, consequentemente, levando a uma redução na produção de forragem e, na capacidade de animais por área.

O controle de invasoras pode ocorrer de forma preventiva que consiste na adoção de práticas para prevenir a introdução, estabelecimento ou a disseminação de determinadas espécies daninhas em áreas ainda não infestadas. Uma segunda estratégia é o controle mecânico (físico), onde é feita o combate manual ou automatizada da planta, normalmente através das roçadas, que podem não ser suficientes para o controle e eliminação das invasoras. Por fim, tem-se o controle químico, que envolve a aplicação de herbicidas para eliminar a planta daninha. No entanto, o produto aplicado deve ser totalmente seletivo à forrageira, não influenciando seu desenvolvimento fenológico ou prejudicando o seu rendimento. Os autores destacam ainda que, dada a diversidade de possíveis espécies presentes nas pastagens, pode ser necessária a aplicação de misturas de herbicidas (Pereira et al., 2011).

O controle químico tradicional de plantas daninhas, pode sobrecarregar os custos da produção pecuária, pois sem a exata localização das plantas invasoras, certa quantidade herbicidas é desperdiçada uma vez que costuma se fazer aplicação homogênea nas áreas. Além de que, alguns herbicidas são prejudiciais ao meio ambiente. Uma das possíveis alternativas com relação ao problema, seria

a automatização da identificação das plantas invasoras e seu mapeamento para uma aplicação mais adequada dos insumos.

A proposta desta pesquisa é desenvolver estratégias automáticas que possam auxiliar esse processo de identificação e combate às ervas daninhas. A premissa é de que a partir de imagens de regiões de pastagem seja possível submetê-las a técnicas de processamento de imagens combinadas com algoritmos de aprendizagem de máquina de forma a identificar elementos não pertencentes às pastagens e, dentre estes, detectar quais são possíveis plantas invasoras para a aplicação efetiva de herbicidas. Esse processo pode contribuir para agilizar o tratamento, bem como acarretar economia para o produtor.

Métodos convencionais de programação geralmente não conseguem reconhecer e classificar essas plantas com a precisão necessária. Por isso, é essencial recorrer a técnicas de análise de imagens e de aprendizagem de máquina, que conseguem identificar padrões em conjuntos que envolvam informações variadas. A aprendizagem de máquina oferece uma abordagem mais flexível, já que pode se adaptar a diferentes conjuntos de dados e condições, sem a necessidade de reescrita de todo o algoritmo. Com um processo bem planejado de extração de características e seleção de algoritmos de classificação, pequenas alterações nos dados podem ser solucionadas ajustando-se apenas a etapa de calibração dos modelos de aprendizagem. Isso proporciona maior robustez ao projeto e permite uma evolução contínua do sistema, mantendo sua relevância e eficácia à medida que as condições do mundo real mudam.

Tendo em vista este problema busca-se, neste trabalho, propor e avaliar um protocolo automático capaz de, a partir de imagens de pastagens (com o sem a presença de plantas invasoras), extrair características discriminantes e treinar modelos de aprendizagem de máquina que sejam capazes de identificar a presença de ervas daninhas. Espera-se que a solução proposta seja capaz de classificar com precisão regiões de forrageiras e áreas contendo plantas invasoras permitindo a aplicação pontual de herbicidas, diminuindo assim o investimento com produtos químicos, mão de obra especializada e possíveis danos ao ambiente.

2 Fundamentação

Segundo Faceli et al. (2011), a resolução de problemas computacionais tradicionalmente envolve a escrita de algoritmos detalhados, que especificam passo a passo como solucionar uma questão. No entanto, para certas tarefas, como o reconhecimento facial ou de voz, esses passos não são tão fáceis de se definir. Nós, seres humanos, conseguimos identificar rostos e vozes de maneira intuitiva, mas transformar essa habilidade em código é uma tarefa complexa.

Em contextos como o reconhecimento de plantas daninhas em pastagens, a abordagem tradicional de algoritmos pode ser inadequada. Pequenas variações nas condições, como iluminação deficiente ou variação de características, podem comprometer a eficácia do algoritmo. Assim, métodos tradicionais podem ser ineficazes para problemas que exigem reconhecimento de padrões complexos. Neste

sentido a aprendizagem de máquina pode ser uma abordagem alternativa que pode identificar padrões e adaptar-se a dados variáveis, tornando-se uma solução mais robusta.

Nesse contexto, para a realização do processo de classificação aqui proposto foram avaliadas diversas estratégias de aprendizado supervisionado, como K-Nearest Neighbors (Aha et al., 1991), que utiliza a classe dos vizinhos mais próximo na classificação de uma instância. Neste algoritmo, o espaço é definido pelo conjunto de atributos de forma que os exemplares do conjunto estão distribuídos nesse espaço. O algoritmo, ao receber um novo exemplar, cuja classe é desconhecida, encontra seus k vizinhos mais próximos nesse espaço de características. Em seguida ele calcula a classe mais frequente entre os vizinhos escolhidos, atribuindo esta ao novo exemplar.

O *Support Vector Machines* (SVM), proposto por Platt (1998), segue a ideia da construção de um hiperplano ótimo para a separação das classes. Caso as classes não sejam linearmente separáveis o método pode aplicar uma transformação nos dados, elevando o espaço de características (*kernel trick*). Esse processo pode ser repetido até que seja obtido um hiperplano capaz de separar as diferentes classes ou que um erro máximo seja obtido.

Outra abordagem utilizada são as redes neurais artificiais, que se baseiam no princípio da conexão neural humana para fazer a troca de informações. Um exemplar bastante utilizado neste contexto é o Perceptron de Múltiplas Camadas (*Multilayer Perceptron* - MLP). Sua estrutura genérica é composta por Perceptrons (análogos aos neurônios) distribuídos na camada de entrada, camadas escondidas e camada de saída. Cada perceptron é capaz de construir apenas um modelo linear de classificação. No entanto, ao agruparmos vários exemplares, como no MLP, o modelo é capaz de lidar com problemas mais complexos. O aprendizado consiste na calibração do peso atribuído às conexões dos neurônios presentes na rede.

Uma árvore de decisão (*Decision Tree* - DT) pode ser entendida como grafo acíclico direcionado (representado através de uma estrutura hierárquica), onde cada nó interno representa uma decisão (ou teste) sobre um atributo, cada ramo corresponde ao resultado do teste, e cada nó folha indica uma classe ou valor de saída (Castro and Ferrari, 2016). O objetivo é criar uma árvore que particione o espaço de dados em regiões que correspondam a diferentes classes ou resultados.

A Floresta Aleatória (RF - *Random Forest*), proposto por Breiman (2001), é um meta estimador que ajusta múltiplos classificadores de árvores de decisão em subamostras do conjunto de dados, usando a média (ou moda) das previsões para aumentar a precisão e reduzir o *overfitting* (Pedregosa et al., 2011). Portanto, as previsões das várias árvores são analisadas em conjunto, buscando uma maior precisão dos resultados obtidos. Normalmente esta estratégia se sobressai às DTs, uma vez que apresenta maior robustez e eventualmente levando a melhores taxas de acerto.

Técnicas específicas de aprendizagem de máquina podem obter um melhor resultado dependendo do problema ou da maneira a qual ela foi empregada. Por este motivo ao utilizar tais técnicas e algoritmos, testes devem ser feitos, para a verificação de qual método mais se ajusta ao problema.

Porém, em problemas amplos, como os que envolvem reconhecimento de imagens, diferentes técnicas podem apresentar uma eficácia maior em diferentes imagens. Visando aproveitar os pontos fortes de cada método de aprendizagem, pode-se utilizar estratégias baseadas em *ensembles*, os Sistemas de Múltiplos Classificadores (SMC) que, segundo alguma estratégia estabelecido, combinam vários modelos para se obter um valor de saída.

Um SMC parte de pressuposto de que em problemas de maior variabilidade, uma solução monolítica (baseada em um único modelo de classificação), pode não ser robusta suficiente. O pressuposto é de que os classificadores em um SMC são diversos no sentido de que geralmente cometem erros distintos e, dessa forma, podem ser complementares (Araújo et al., 2017).

As estratégias mais comuns de classificação, conforme apresentado por Gunes et al. (2003), são apresentadas a seguir.

- Voto majoritário: Cada modelo dá sua opinião e a classe que receber mais votos entre os 5 classificadores, será considerada como resposta da estratégia.
- Regra da soma: A classe que tiver um somatório maior da confiança dos votos de cada modelo, é considerada a saída do SMC.
- Regra da multiplicação: A multiplicação da confiança de cada classificador é feita, e a classe com o maior resultado é considerada.
- Regra do máximo: A classe que obtiver a maior confiança de algum dos classificadores é tida como resposta.
- Regra da mediana: O cálculo da mediana é feito sobre a confiança dos métodos sobre as 2 classes em questão, a classe com a maior mediana é considerada.

Outras estratégias podem, no entanto, serem aplicadas, como as combinações por ranking ou Borda Count, na qual as classes recebem pesos de acordo com o grau de confiança que os classificadores têm em seus votos.

2.1 Trabalhos Correlatos

Na literatura, diversas pesquisas utilizam técnicas de aprendizagem de máquina e inteligência artificial para a detecção de plantas prejudiciais a pastagem ou a plantações. Neste contexto, pode citar a obra de Smith et al. (2019) que visa a identificação de daninhas, como a classe *Rumex* e o trevo branco em meio ao pasto, utilizando Redes Neurais Convolucionais (*Convolutional Neural Networks* - CNNs). Os experimentos foram divididos em dois cenários. No primeiro comparou-se apenas a classe pastagem com a *Rumex* e no segundo foi engloba também a classe trevo branco. No trabalho os autores aplicaram uma estrutura baseada na ResNET (*Residual Network*) frente a estratégias baseadas em transferência de aprendizado. Os resultados mostraram que a rede proposta foi capaz de acertar 95,6% das instâncias no primeiro cenário e 94,9% no segundo.

Em seu trabalho, Zhang et al. (2018) propuseram uma estratégia automática para detecção de ervas daninhas de folha larga em pastagens baseada em redes neurais convolucionais. Os autores compararam o desempenho de estratégias de aprendizagem de máquina convencionais como SVM, KNN frente à estratégia proposta. Além da

estratégia de CNN foi realizada a captura de imagens e construção de um conjunto composto por 6087 exemplares, dos quais 4080 continham apenas pastagem e 2007 continham a presença da planta invasora. Os resultados observados mostraram que o SVM pode obter acurácia de 89,4%, o KNN alcançou uma taxa de acertos próxima de 85% enquanto a estratégia proposta atingiu acurácia de 96,88%.

Esposito et al. (2021) apresentam um resumo do avanço da agricultura com utilização de novas tecnologias, com ênfase no uso de aeronave remotamente pilotada (RPA – *Remotely Piloted Aircraft*) e na aplicação de sensores RGB e hiperespectrais de diversos modelos. Os autores descrevem as etapas necessárias e suas importâncias, na utilização de sensores para a identificação de plantas invasoras. Além disso, destacam, que relevantes conhecimentos têm sido gerados em consequência da análise de imagens remotas neste contexto, como por exemplo, a redução da perda de colheitas de trigo no inverno, em função da presença de uma diversidade de daninhas. O estudo conclui que para ampliar a abordagem apresentada em contextos reais e eliminar especificamente apenas daninhas prejudiciais ao plantio, novas percepções sobre a dinâmica populacional de ervas daninhas devem ser geradas.

No trabalho de Li et al. (2021), feito na Nova Zelândia, foram utilizadas imagens hiperespectrais de duas espécies de grama (*Setaria pumilla* e *Stipa arundinacea*) e duas espécies de plantas daninhas de folha larga (*Ranunculus acris* e *Cirsium arvense*). Foram adotadas duas abordagens para o processamento das imagens. Na primeira, as imagens foram divididas em subconjuntos removendo-se os subconjuntos não pertencentes à planta, e então levantando-se a média dos valores de matiz, saturação e intensidade (HSI – *Hue, Saturation, Intensity*) de cada subconjunto. Na segunda estratégia não houve divisão em subconjuntos, mas o fundo da imagem foi removido por meio de uma máscara e, em seguida, a média dos valores HSI foi calculada.

Além disso, foram treinados três modelos de classificação: análise discriminante por mínimos quadrados parciais, SVM e MLP. Todos os modelos de classificação conseguiram identificar as quatro classes de plantas propostas, com a acurácia variando entre 70% e 100%. Entretanto, o modelo mais confiável e robusto foi o MLP, alcançando uma taxa de acerto de 89,1%. Ademais, os autores concluem que modelos desenvolvidos com superpixels podem prover resultados melhores do que simplesmente realizar a média sobre a imagem inteira.

Em sua pesquisa, Espejo-Garcia et al. (2021) empregaram duas diferentes bases de dados em conjunto, a primeira é composta por 504 imagens de quatro diferentes espécies de plantas nos primeiros estágios de crescimento, enquanto a segunda é formada por 960 imagens de 12 diferentes espécies, em diversas etapas do crescimento. Além disso, foi também utilizado *Automated Machine Learning* (AutoML), o qual foi responsável por extrair os melhores atributos das imagens, reduzir a dimensão e realizar os testes. Para esse processo, diversas técnicas foram utilizadas, as principais foram; *Singular Value Decomposition* (SVD), ANOVA, *Decision Tree*, *AdaBoosting*, *Ensemble*, *Single*, *Softmax*, *Random Forests*, *Extra Trees* e *Bayesian optimization*. A métrica empregada como critério avaliativo foi a F1-Score. Os resultados obtidos no estudo foram promissores,

com valores de F1-score de 93,8% e 90,74% dependendo da base de dados utilizada. Além disso, os autores realizaram experimentos adicionando ruídos e desfoque nas imagens visando avaliar a robustez da solução proposta. Os resultados mostraram-se um tanto inconclusivos, com alguns cenários em que houve melhora no índice F1-Score, mas houve casos em que se observou piora de até 5,23%.

Outro trabalho importante é o de Wu et al. (2021), em que foi realizada uma revisão de estratégias para detecção de ervas daninhas em diferentes contextos. Os autores detalham os trabalhos voltados à área, identificando as características das imagens, modelos empregados, características utilizadas na detecção e resultados alcançados. As análises indicaram que a literatura ainda carecia de conjuntos de dados públicos que sirvam de *baseline* para outras pesquisas. Além disso, os pesquisadores indicaram que as tarefas de detecção e identificação são muito sensíveis às variações de características de diferentes ervas daninhas, de propriedades das imagens colhidas, de cenários onde há a sobreposição de espécies e mesmo fatores como custo computacional para o processamento. Segundo os autores, apesar das boas taxas de sucesso a aplicação prática das técnicas ainda esbarra nos requisitos necessários para atender às restrições dos modelos, indicando a carência de equipamentos de alto desempenho com custos acessíveis.

No estudo de Sapkota et al. (2022), destaca-se a necessidade de grandes volumes de dados para a utilização de estratégias de *deep learning*. Os autores propuseram uma abordagem alternativa, na qual imagens obtidas por meio de RPAs foram cortadas e processadas para gerar imagens sintéticas através de uma rede adversária generativa (GAN – *Generative Adversarial Network*). Essas imagens foram, então, utilizadas para treinar uma rede neural convolucional do tipo Mask R-CNN, visando a detecção de plantas invasoras. O corte das imagens foi realizado tanto de maneira manual quanto automática, permitindo uma comparação entre os métodos. As sub-imagens foram submetidas ao processo de *data augmentation* onde foram aplicadas operações aleatórias de rotação, mudança de cor e de escala. Ademais, foram coletadas 99 amostras de biomassa da *Ipomoea spp.* e 104 amostras de grama. Para a predição da biomassa, foram utilizados 2 metodologias: *Canopy Mask* e *Bounding Box*.

Os autores puderam concluir que imagens sintéticas podem ser uma boa alternativa, pois alcançaram aproximadamente 80% da acurácia obtida com imagens reais. As bases de dados composta por imagens cortadas automaticamente, proveram resultados muito próximos da base produzida com imagens cortadas manualmente, mostrando que tempo e outros recursos podem ser poupados utilizando esta abordagem. Sobre a predição da biomassa, a metodologia *Canopy Mask* demonstrou desempenho superior ao *Bounding Box*. A predição de biomassa é importante pois oferece a oportunidade de desenvolver um regime de manejo de pastagem específico do local para melhorar a lucratividade.

3 Metodologia

As etapas implementadas nesta pesquisa são ilustradas na Fig. 1. No primeiro passo (A), cada imagem proveniente do conjunto de entrada é submetida à amostragem

de exemplares da classe negativa, os quais correspondem às amostras de pastagem (quadro superior destacados em vermelho) e da classe positiva (quadro inferior) onde tem-se as amostras da planta invasora. Cada amostra é então submetida ao processo de extração de características (B) de forma que cada uma será armazenada como um registro estruturado em um conjunto de dados.

O conjunto de dados é originalmente dividido em conjunto de treino, que é empregado no aprendizado dos modelos (etapa C), validação, utilizado para a calibração de hiperparâmetros dos modelos (etapa D) e, por fim, conjunto de teste, o qual serve de base para a avaliação dos desempenhos dos classificadores (etapa E). Cada uma das etapas será mais bem detalhada nas seções seguintes.

Na implementação, que se deu em linguagem Python, utilizou-se as bibliotecas Pandas para manipulação de dados, Matplotlib e Seaborn para exibição de informações de forma gráfica, bem como a biblioteca Sklearn para o desenvolvimento das estratégias de aprendizado de máquina. Além disso, foi empregada a biblioteca OpenCV para o processamento das imagens.

3.1 Base de Dados

O processamento das imagens, deve ser feito a partir de exemplares obtidos com o uso de câmeras ou ARPs. Algoritmos avançados analisam essas imagens para detectar plantas daninhas, diferenciando-as da vegetação forrageira.

O conjunto de imagens utilizado, provido por [Terada \(2022\)](#), é formado por 1114 imagens de pastagens que podem ou não conter a presença de uma planta invasora que, neste caso, é representada pelo gênero *Rumex*. Um dos exemplares mais comuns do gênero em terras brasileiras, principalmente no Sul do país, é a *Rumex obtusifolius* (Língua-de-vaca), e que, segundo [Jochem et al. \(2020\)](#), pode ser encontrada em plantações, beiras de estradas, pastagens e até mesmo em áreas urbanas.

A Língua-de-vaca, que possui caráter perene, pode se reproduzir por meio de sementes e rizomas, o que dificulta o seu combate ([Güttler et al., 2018](#)). Ela apresenta alto potencial competitivo visto que pode florescer duas vezes ao ano e por produzir um número elevado de sementes (cerca de 60.000 por planta) as quais podem permanecer viáveis por vários anos ([Jochem et al., 2020](#)).

Além da alta capacidade de reprodução, [Zaller \(2006\)](#) aponta que a planta pode rebrotar após o corte de sua parte aérea (como feito durante a roçada). Ademais, o autor indica que a Língua-de-vaca tem efeito alelopático sobre cinco espécies de gramíneas, podendo inibir a sua germinação.

O conjunto de fotos utilizadas, provido por [Terada \(2022\)](#), é formado por 1114 imagens de pastagens que são divididas em três subconjuntos. O primeiro, treino, contém 763 exemplares. Já o segundo (validação) e terceiro (teste) são compostos, respectivamente, por 228 e 123 imagens.

Cada imagem presente na base de dados, no entanto, pode conter várias amostras de *Rumex*, como é ilustrado na [Fig. 2](#). Pode-se perceber que há sete sub-regiões nas quais estão destacados exemplares da planta invasora. Tal demarcação foi feita pelos autores da base, de forma que

as posições de cada sub-região já estão mapeadas.

As imagens fornecidas têm dimensão 640 x 640 pixels. As sub-regiões contendo as amostras de *Rumex*, no entanto, têm tamanho variado, de acordo com a dimensão da erva à qual o recorte está destacando. Tal fato é evidenciado na [Fig. 2](#) onde pode-se perceber que as sete regiões destacadas têm extensões variadas.

Uma vez que cada imagem presente no conjunto pode conter diversas amostras da *Rumex*, fez-se o levantamento de quantos exemplares da planta estavam contidos em cada conjunto. Nas imagens pertencentes ao grupo de teste estão presentes 572 amostras positivas para a erva. O conjunto de validação, por sua vez, contém 805 exemplares. Já no conjunto de treino estão presentes 2460 amostras da classe positiva (planta invasora).

O conjunto fornecido, entretanto, não contém casos negativos para a erva daninha, ou seja, as subamostras evidenciam apenas os casos em que há a ocorrência da *Rumex* (que é a classe positiva), não contendo exemplares de pastagem. Foi necessária então, a criação de conjuntos representativos de imagens de pastagem em que não houvesse a presença da planta invasora para formar o conjunto de amostras para a classe negativa.

3.2 Obtenção das amostras

Uma vez que as amostras de casos positivos já haviam sido determinadas pelos autores da base de dados, foi realizado o processo de construção de dois conjuntos de amostras de sub-imagens contendo apenas pastagem. Essa demanda dá-se pelo fato de que os modelos de aprendizagem de máquina têm a necessidade de “conhecer” exemplares das duas classes para que possa aprender a diferenciá-las.

Visando formar conjuntos representativos esses recortes foram determinados de forma aleatória, ou seja, a posição e dimensão da amostra negativa eram determinadas randomicamente. No entanto, duas restrições precisavam ser respeitadas: A primeira era de que a amostra de pastagem não poderia ter sobreposição com os recortes das classes positivas; e a segunda restrição tem relação com tamanho das amostras. A dimensão de uma sub-imagem variou entre a média dos tamanhos dos recortes de erva daninha variadas de um patamar fixo. O valor empregado para tal foi de 10 pixels, tanto para a altura quanto para a largura da imagem. O limite foi definido empiricamente, e permitiu que as amostras não contivessem regiões muito pequenas ou demasiadamente grandes.

A [Fig. 3](#) ilustra um cenário das amostras obtidas a partir do processo descrito. Os recortes na cor roxa representam amostras da classe positiva, indicando a presença da planta invasora. Estas marcações foram realizadas pelos autores do conjunto de dados. Já na cor laranja, estão destacadas as regiões recortadas para a formação de amostras da classe negativa, contendo apenas pastagem.

Para cada imagem da base original, foram separadas todas as plantas daninhas, além de algumas amostras de grama. A quantidade de amostras de pastagem utilizadas foi calculada com base no número de imagens de plantas daninhas extraídas de cada imagem. Para avaliação dos métodos aplicados foram construídos dois conjuntos de exemplares da classe negativa. O primeiro deles, nomeado “1.0” contém a mesma quantidade de amostras

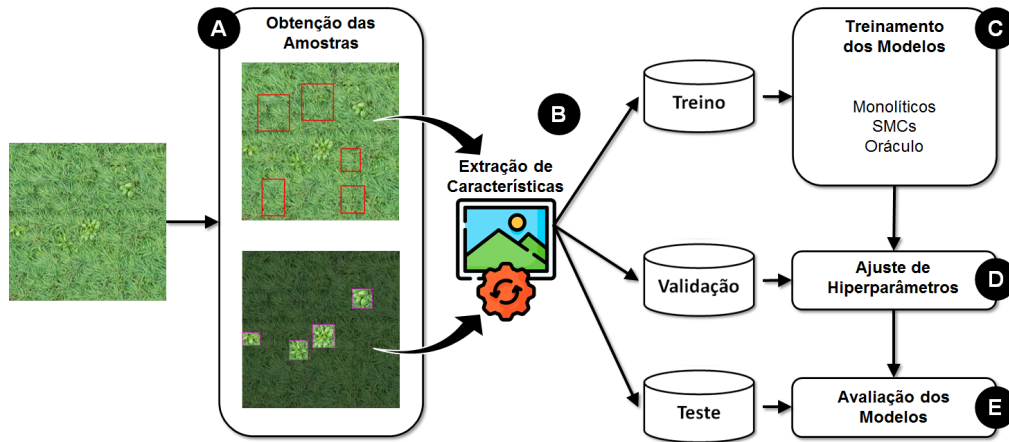


Figura 1: Fluxograma contendo as etapas desenvolvidas e como cada uma se comunica com as demais



Figura 2: Exemplo de amostra presente no conjunto de dados empregado nos experimentos. Em destaque sete ervas daninhas presentes na captura
Fonte: [Terada \(2022\)](#)

positivas e negativas, ou seja, a base tem uma distribuição exatamente balanceada entre as classes. No entanto, em um cenário real, será mais comum ter amostras negativas, visto que a pastagem é muito mais presente do que a planta invasora. Dessa forma, construiu-se um segundo conjunto de dados, chamado “2.0”, em que o número de amostras de pastagem é duas vezes maior do que o número de registros contendo imagens da *Rumex*.

A distribuição de instâncias entre os dois conjuntos de dados é detalhada na [Tabela 1](#) onde são apresentadas as quantidades de exemplares de cada classe em cada um dos três conjuntos (treino, teste e validação).



Figura 3: Resultados dos recortes realizados. Os retângulos em roxo representam as plantas invasoras presentes na imagem, enquanto aqueles na cor laranja, os exemplares recortados a partir de localizações aleatórias da imagem, representando a classe pastagem

3.3 Extração de características

Cada recorte realizado, como demonstrado na [Fig. 3](#) tornou-se um registro no conjunto de dados. Para a utilização dos modelos de aprendizagem de máquina, no entanto, essas sub-imagens foram convertidas para conjuntos de dados estruturados. Nesse processo foram extraídos os valores médios dos canais R (Red), G (Green), H (Hue), S (Saturation) e V (Value) do grupo de pixels presentes no recorte. Além disso, foram levantados os desvios padrões dos seis canais citados, obtendo-se assim um registro contendo 12 atributos descritores e uma 13ª coluna representando a classe da sub-imagem.

A escolha de fornecer aos métodos de aprendizagem, o desvio padrão referentes aos pixels da imagem teve como motivação o fato de que em uma imagem de amostra posi-

tiva e outra negativa, a média dos valores dos canais pode ser a mesma, porém o desvio padrão provavelmente vai ser diferente, devido ao contraste existente entre as plantas daninhas e a pastagem.

Na Fig. 4, são apresentados os histogramas dos canais RGB e HSV para um exemplar da classe positiva (representado na coluna direita da figura) e outro da classe negativa (representado na coluna esquerda). As duas imagens analisadas tinham a mesma dimensão, 133 x 152 pixels.

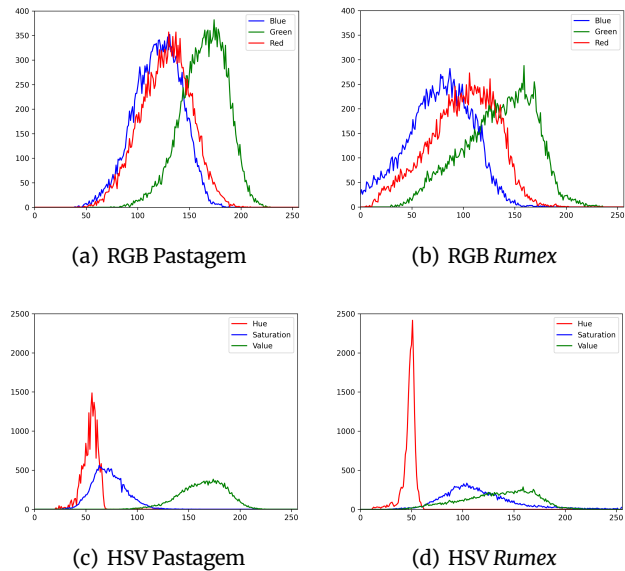


Figura 4: Histogramas dos canais RGB e HSV para exemplares da forrageira (lado esquerdo) e da planta invasora (lado direito). As Fig. 4(a) e Fig. 4(b) apresentam a contagem de pixels dos canais R (em vermelho), G (em verde) e B (em azul). Já as Fig. 4(c) e Fig. 4(d) apresentam a contagem de pixels dos canais H (em vermelho), S (em azul) e V (em verde)

Ao analisarmos os histogramas dos canais RGB apresentados nas Fig. 4(a) e Fig. 4(b), percebemos algumas diferenças entre o comportamento apresentado pela planta invasora e a pastagem. O primeiro fato observado é que os tons que compõem a Rumex são mais escuros, o que é evidenciado pelo maior acúmulo de frequência no lado esquerdo do histograma. Além disso, a sua distribuição ao longo do espaço é mais abrangente em comparação ao observado no histograma das pastagens, onde os canais têm picos maiores de frequência, mas estão mais concentrados

em regiões menores do espaço de cores. Pode-se perceber também que os canais R e G da imagem referente à forrageira têm mais sobreposição do que os mesmos canais para a planta invasora.

Com relação aos canais H, S e V, percebe-se que o canal Hue (em vermelho) quase não tem sobreposição com os outros dois no caso da Rumex. Já para a pastagem, o canal tem sobreposição com a saturação (em azul). Além disso, nota-se que a distribuição dos três canais para a grama é mais em determinadas regiões do histograma enquanto para a planta invasora a frequência dos valores é mais distribuída. O exemplo evidencia ainda que os canais Saturation e Value (em verde) quase não têm sobreposição nos histogramas referentes à forrageira, fato que não é constatado no histograma da erva daninha.

Outra vantagem de empregar os descritores citados é de que eles são invariantes a escala, deslocamento e rotação.

3.4 Treinamento e Ajuste de Hiperparâmetros

A etapa de aprendizado e calibração dos modelos baseou-se em dois conjuntos de dados. O conjunto de treino foi usado durante a construção dos modelos e as amostras de validação foram empregadas durante o ajuste da melhor configuração dos modelos construídos.

O treinamento adequado dos modelos e seu desempenho dependem diretamente de uma boa definição de hiperparâmetros. Visando obter a configuração mais promissora para nossos experimentos, implementamos uma estratégia baseada em *grid search* para definir qual configuração deveria ser adotada. O critério utilizado para orientar nossa escolha foi a acurácia, que indica a porcentagem total de exemplos corretamente classificados pelo modelo.

Visando evitar *overfitting* ou vieses, a escolha baseada em grade é feita sobre o conjunto de validação. Assim, treinamos cada um dos cinco modelos com todas as combinações possíveis dos valores de hiperparâmetros e selecionamos aquelas que levaram à maior acurácia do classificador sobre o conjunto de validação. Os valores avaliados nesta etapa são apresentados na Tabela 2. Nela são apresentados todos os hiperparâmetros que foram variados (segunda coluna), bem como os valores testados (terceira coluna) e aqueles que levaram à melhor configuração. Os valores apresentados nas colunas de melhor configuração são os mais frequentemente escolhidos ao longo das 20 execuções dos modelos.

4 Resultados e Discussão

Após identificadas as melhores configurações para cada classificador, estes foram aplicados sobre os conjuntos de

Tabela 1: Distribuição do número de amostras para cada conjunto de dados

Dataset	Conjunto 1.0		Conjunto 2.0	
	Quantidade de Amostras Positivas	Quantidade de Amostras Negativas	Quantidade de Amostras Positivas	Quantidade de Amostras Negativas
Treino	2460	2460	2460	4920
Validação	805	805	805	1610
Teste	572	572	572	1144
Total	3837	3837	3837	7674

testes com o objetivo de avaliar seus desempenhos. Os resultados são compostos por medições baseadas na acurácia, sensibilidade e especificidade dos modelos.

A sensibilidade corresponde à porcentagem de acertos entre os exemplos positivos, representados pelas plantas invasoras no conjunto de dados, ou seja, ela reflete a capacidade dos modelos em identificar corretamente a presença das ervas daninhas. Já a especificidade reflete a porcentagem de acertos entre os exemplos negativos que, neste caso, correspondem à grama.

Um critério que pode ser utilizado para entender o desempenho de um ensemble de classificadores é o oráculo. Ele é o limite máximo de acertos que qualquer combinação dos classificadores presentes no conjunto pode alcançar. Para se determinar o desempenho do oráculo, para cada instância presente no conjunto de teste é verificado se algum dos classificadores conseguiu acertar a sua classificação. Então o oráculo considera que o SMC é capaz de escolher aqueles ou aqueles classificadores capazes de acertar aquela instância.

Para ilustrar o comportamento dos modelos de classificação selecionou-se uma imagem do conjunto de entrada para que o KNN, que foi aquele que obteve a maior acurácia entre todos os experimentos, fizesse a classificação dos recortes presentes na imagem. Tal processo é apresentado na Fig. 5. Na Fig. 5(a) é exibida a imagem de entrada contendo 16 amostras, 8 pertencentes à classe negativa (em laranja) e 8 à classe positiva (em roxo). Na Fig. 5(b) é exibida a mesma imagem de entrada tendo a saída da opinião do classificador KNN. Os recortes destacados na cor verde indicam as amostras classificadas corretamente enquanto as marcações em vermelho indicam erros cometidos pelo classificador. Pode-se perceber que ele não conseguiu reconhecer 2 das 16 sub-imagens fornecidas, tais amostras deveriam ser classificadas como plantas invasoras, porém o classificador não foi capaz de chegar a esta conclusão.

Na Tabela 3, é apresentada a acurácia média (e seu respectivo desvio padrão) de cada modelo que foi calculada ao longo de 20 execuções sobre o conjunto de teste para as duas bases de dados representadas na Tabela 1. Nela, podemos verificar quais modelos apresentam melhor desempenho no reconhecimento dos exemplares fornecidos no conjunto de teste. Além dos cinco modelos de classi-

ficação apresentados na Tabela 2 foram avaliadas cinco estratégias de combinação de classificadores: as regras da soma, produto, média e máximo, além da combinação pelo voto majoritário.

Tabela 3: Acurácia média dos modelos sobre o conjunto de teste ao longo de 20 execuções. Entre parênteses são apresentados os desvios-padrões

Estratégia de Classificação	Acurácia Conjunto 1.0	Acurácia Conjunto 2.0
KNN	0,9519 (0,00000)	0,9592 (0,00000)
AD	0,9410 (0,00043)	0,9441 (0,00058)
SVM	0,9249 (0,00019)	0,9242 (0,00052)
RF	0,9552 (0,00183)	0,9517 (0,00197)
MLP	0,9444 (0,00893)	0,9505 (0,00818)
Voto Majoritário	0,9524 (0,00292)	0,9518 (0,00155)
Regra da Soma	0,9538 (0,00276)	0,9510 (0,00144)
Regra do Produto	0,9555 (0,00203)	0,9507 (0,00190)
Regra da Mediana	0,9541 (0,00311)	0,9501 (0,00247)
Regra do Máximo	0,9545 (0,00330)	0,9503 (0,00238)
Oráculo	0,9765 (0,00115)	0,9768 (0,00165)

Analizando-se os dados apresentados na Tabela 3, nota-se que os modelos com os melhores desempenhos para o primeiro conjunto de dados foi a combinação de todos os modelos baseando-se na regra da multiplicação (95,55%), seguida pelo Random Forest, com 95,52% de acertos e pela combinação pela regra do máximo, que obteve acurácia de 95,45%. As piores taxas de acerto foram obtidas pela árvore de decisão (94,10%) e SVM, com 92,49% de acertos.

Para o segundo conjunto de dados, em que o número de amostras negativas possui o dobro de instâncias da classe negativa, quem apresentou o melhor desempenho foi o KNN, acertando 95,92% das amostras de teste. Em seguida, os melhores desempenhos foram obtidos pela combinação pelo voto majoritário e pelo Random Forest, com acurácia de 95,18% e 95,17%, respectivamente.

É notável que independente da base de dados utilizada, os resultados foram similares, ou seja, todos os modelos se comportaram bem, mostrando que a extração dos dados foi feita de forma correta, trazendo atributos úteis, para a

Tabela 2: Conjunto de hiperparâmetros calibrados por meio do Grid-search

Modelo	Hiperparâmetros	Valores Testados	Melhor Config. 1.0	Melhor Config. 2.0
KNN	n_neighbors	[1,2,3..30]	10	19
	weights	[uniform,distance]	uniform	distance
	criterion	[entropy, gini]	entropy	gini
AD	max_depth	[1,2,3..11]	5	9
	min_samples_split	[2,3,4..16]	7	2
	min_samples_leaf	[1,2,3..16]	1	11
SVM	kernel	[rbf,sigmoid,linear,poly]	linear	linear
	C	[0.1,0.2,...,1.0]	0.8	0.9
RF	criterion	[entropy, gini]	entropy	gini
	max_depth	[3,4,5..10]	9	9
	n_estimators	[50,100,150,200]	150	100
	min_samples_leaf	[1,2,3,4,5]	1	3
MLP	hidden_layer_sizes	[(10,10,10),(11,11,11)..(30,30,30)]	(24,24,24)	(16, 16, 16)
	activation	[identity, logistic, tanh, relu]	logistic	logistic
	max_iter	[50,100,150,300,500,1000,10000]	150	1000
	learning_rate	[constant, invscaling, adaptive]	invscaling	invscaling



Figura 5: Ao lado esquerdo, os cortes e as sub-imagens obtidas (em roxo a classe positiva e em laranja a classe negativa). Na direita, quais exemplares o classificador acertou (em verde) e quais errou (em vermelho)

classificação dos segmentos de imagens.

Ao compararmos os desempenhos de todos as estratégias de classificação sobre os dois conjuntos de dados (1.0 e 2.0), percebemos que três abordagens apresentaram certa melhora: o KNN, que teve aumento de 0,73 pontos percentuais na acurácia, a árvore de decisão que evoluiu de uma taxa de acerto de 94,10% para 94,41% e o MLP que apresentou uma evolução de 0,61 pontos percentuais na acurácia. Em todos os outros casos, no entanto, observou-se uma pequena variação negativa na proporção de instâncias classificadas corretamente.

Ao testar o conjunto de acurácias obtidos das 20 execuções para o conjunto de dados 1.0, utilizando o Kruskal-Wallis com 5% de significância, rejeitou-se a hipótese nula ($p\text{-value} = 1,59248E-21$), de que os classificadores não apresentaram desempenhos significativamente diferentes. Tal fato indicou que pelo menos um dos classificadores desempenhou diferente dos demais. A aplicação do teste de Kruskal-Wallis sobre o conjunto 2.0, sob o mesmo grau de confiança, indicou um $p\text{-value} = 7,45391E-23$, rejeitando-se assim a hipótese nula.

Uma vez que para os dois conjuntos H_0 foi rejeitada, aplicamos o teste U de Mann-Whitney (com 95% de confiança) para a comparação das acurácias dos classificadores, analisando-as par-a-par. Os valores obtidos para o $p\text{-value}$ de cada comparação são apresentados na Tabela 4 para o conjunto 1.0 e na Tabela 5 para o conjunto 2.0. Em negrito estão representados os casos que se observou uma diferença significativa.

Visando entender a capacidade dos modelos em discernir as classes específicas foram levantados os índices de sensibilidade e especificidade. Na Fig. 6 são apresentados os valores médios de sensibilidade para as 20 execuções de cada classificador e na Fig. 7 compara-se, visualmente, os valores médios da especificidade dos cinco classificadores ao longo das vinte execuções.

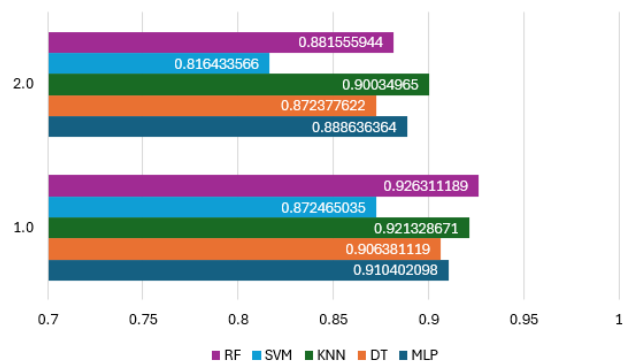


Figura 6: Média do índice de sensibilidade dos classificadores ao longo das 20 execuções

Como esperado, os valores obtidos para a sensibilidade foram maiores para o conjunto 1.0. Tal fato dá-se pela distribuição semelhante entre os exemplares da classe positiva e classe negativa. Ao utilizarmos o conjunto 2.0 percebe-se que todos os modelos tiveram queda na sua capacidade de classificar corretamente a classe positiva. Como os modelos tendem a acertar mais a classe mais frequente e este conjunto contém o dobro de exemplares da classe negativa em comparação à positiva, era esperado que houvesse esta diminuição.

Como pode ser visto na Fig. 6, a melhor sensibilidade entre as duas bases de dados, foi alcançada pelo classificador Random Forest, obtendo 92,6% de acerto sobre os exemplares positivos presentes na base de dados 1.0. Ou seja, este classificador com as configurações encontradas, possui um melhor desempenho na classificação de plantas daninhas, sem levar em conta sua capacidade de

Tabela 4: *p-values* obtidos a partir do Teste de de Mann Whitney sobre o conjunto 1.0

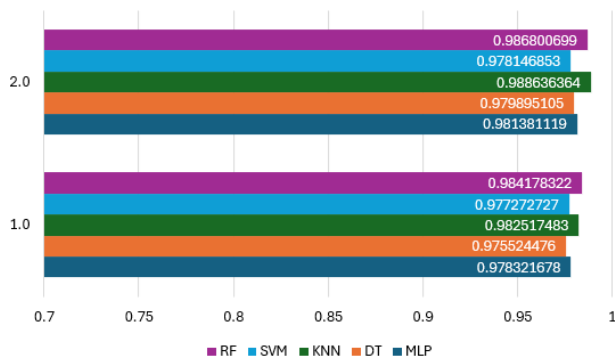
Modelo	MLP	REG_MAX	REG_MEDIANA	REG_SOMA	MAJOR	RF	SVM	DT
KNN	0,00884	0,00015	0,00133	4,17E-02	0,56298	2,78E-08	7,43E-10	4,19E-09
DT	0,12962	4,06E-08	3,97E-08	3,94E-08	3,99E-08	3,89E-08	6,05E-09	-
SVM	2,76E-07	1,10E-08	1,07E-08	1,06E-08	1,08E-08	1,05E-08	-	-
RF	2,05E-05	3,90E-01	0,14765	0,05670	0,00215	-	-	-
MAJOR	0,00340	0,03072	0,07492	0,11747	-	-	-	-
REG_SOMA	0,00039	0,36185	0,57632	-	-	-	-	-
REG_MEDIANA	0,00021	0,69300	-	-	-	-	-	-
REG_MAX	0,00010	-	-	-	-	-	-	-

Tabela 5: *p-values* obtidos a partir do Teste de de Mann Whitney sobre o conjunto 2.0

Modelo	MLP	REG_MAX	REG_MEDIANA	REG_SOMA	MAJOR	RF	SVM	DT
KNN	1,22E-05	7,51E-09	7,83E-09	7,62E-09	7,66E-09	7,73E-09	5,48E-09	4,28E-09
DT	0,00012	3,99E-08	4,13E-08	4,04E-08	4,06E-08	4,09E-08	3,09E-08	-
SVM	7,73E-08	4,83E-08	4,99E-08	4,89E-08	4,91E-08	4,94E-08	-	-
RF	5,60E-01	9,42E-02	0,05909	0,31973	0,84866	-	-	-
MAJOR	0,50626	0,05127	0,04023	1,86E-01	-	-	-	-
REG_SOMA	0,82819	0,30685	0,25320	-	-	-	-	-
REG_MEDIANA	0,49807	0,81704	-	-	-	-	-	-
REG_MAX	0,59642	-	-	-	-	-	-	-

classificação sobre os exemplares negativos. No entanto, ao aumentarmos o número de exemplares da classe negativa o classificador que melhor conseguiu identificar as amostras da classe positiva foi o KNN, com sensibilidade de 0,9003, indicando que ele conseguiu lidar melhor com o desbalanço entre as classes.

Analisando-se a Fig. 7, é perceptível que as taxas de especificidade foram parecidas entre os classificadores, destacando-se novamente o KNN e o Random Forest, com 98,86% e 98,68% de acerto em instâncias negativas, respectivamente, esse desempenho foi alcançado utilizando a base de dados que possui 2 vezes mais exemplares negativos.

**Figura 7:** Média do valor de especificidade dos classificadores durante as 20 execuções

Enquanto para a base de dados com o mesmo número de instâncias positivas e negativas (conjunto 1.0), obteve-se um desempenho um pouco inferior, porém, o resultado alcançado ainda é considerado satisfatório. Novamente liderando a taxa de acertos, temos os classificadores Random Forest e KNN, resultando em 98,41% e 98,25% de acerto para os exemplares negativos.

Uma constatação importante é de que todos os cinco classificadores tiveram melhor desempenho classificando amostras da classe negativa do que identificando os exemplares da planta invasora, tanto para o conjunto balanceado quanto para o cenário em que a classe negativa tinha o dobro de amostras.

Esse fato mostra que a identificação de quais são as plantas daninhas ainda tem margem para melhora. Essa evolução pode se basear em diversos fatores como a extração de outras características que possam discriminar a classe positiva de forma mais contundente, a obtenção de imagens de melhor qualidade ou mesmo a adoção de estratégias de classificação mais avançadas, como redes neurais convolucionais.

A Tabela 6 apresenta uma sintetização dos trabalhos correlatos voltados à tarefa de identificação de plantas invasoras em pastantes. Na tabela são identificadas a aplicação do trabalho assim como as técnicas empregadas e seus desempenhos. Uma consideração pertinente é que uma vez que as bases de dados utilizadas são distintas, não é possível afirmar com precisão qual metodologia obteve o melhor desempenho. Isso se deve ao fato de que cada base apresenta características próprias, como tipos de plantas, qualidade das imagens, distância de captura e métodos de processamento diferentes, o que torna inadequada uma análise comparativa profunda.

Entretanto, é possível observar que, entre os trabalhos correlatos, os melhores resultados foram obtidos por Smith et al. (2019) (que também trabalhou com o gênero *Rumex*), com acurácias de 0,956 e 0,949 utilizando uma rede ResNet. Em comparação, nosso trabalho alcançou uma acurácia de 0,959. Esses resultados não indicam necessariamente que nossa metodologia seja superior, mas permitem concluir que a estrutura proposta é adequada e que os resultados obtidos estão alinhados com os trabalhos de referência na área.

Tabela 6: Comparativo entre os trabalhos correlatos

Trabalho	Tarefa	Métodos Utilizados	Desempenho
Smith et al. (2019)	Identificação de ervas daninhas (<i>Rumex</i>)	ResNet e MobileNet	Acc ResNet = 0,956 e MobileNet = 0,932
Smith et al. (2019)	Identificação de ervas daninhas (<i>Rumex</i> e Trevo Branco)	ResNet e MobileNet	Acc ResNet = 0,949 e MobileNet = 0,911
Zhang et al. (2018)	Detecção de ervas daninhas de folha larga	CNN, SVM e KNN	Acc KNN próximo de 0,85, SVM em torno de 0,894 e CNN de 0,969
Li et al. (2021)	Identificação de duas espécies de gramas e duas ervas daninhas de folha larga usando apenas informações espectrais (Av) ou informações espectrais combinadas com dados espaciais (Sp)	PLS-DA, MLP e SVM	Acc PLS-DA de 0,87 (Av) e 0,7 (Sp), Acc SVM foram 0,92 (Av) e 0,84 (Sp) e Acc MLP foram 1,00 (Av) e 0,89 (Sp)
Espejo-Garcia et al. (2021)	Identificação de ervas daninhas (de pelo menos 12 espécies diferentes)	Singular Value Decomposition, ANOVA, DT, Ada-Boosting, Ensemble, Single, Softmax, RF, Extra Trees e Bayesian optimization	F1 Score Ensemble de 0,938 e 0,9074 para 2 diferentes datasets
Sapkota et al. (2022)	Identificação de plantas invasoras e predição de biomassa utilizando imagens sintéticas	Mask R-CNN	Acc de 0,8

5 Conclusão

A presença de plantas invasoras pode ser prejudicial tanto à vegetação forrageira quanto aos animais pastejando, incorrendo em impactos negativos na produção leiteira ou de corte.

Uma possibilidade de auxiliar o produtor pecuário no combate às ervas daninhas presentes nas pastagens é a detecção automática das invasoras para a aplicação efetiva do combate, seja mecânico ou químico. Dessa forma, neste trabalho realizou-se a aplicação de estratégias de aprendizado de máquina para a detecção de exemplares do gênero *Rumex* (aqui no Brasil representado pela língua-de-vaca).

Foi levantando um conjunto de dados contendo amostras de imagens de pastagens onde os autores identificaram as amostras positivas para a planta invasora. Foi então realizada a construção de um conjunto de dados com exemplares da classe negativa (pastagem) e extração de descritores de texturas, como os canais RGB e HSV. As instâncias foram então submetidas a quatro estratégias monolíticas (KNN, AD, SVM e MLP) e seis abordagens baseadas em sistemas de múltiplos classificadores (RF e combinações pelas regras da soma, produto, mediana, máximo e pelo voto majoritário).

O desempenho observado ao longo dos experimentos indicaram que os atributos extraídos como descritores foram suficientemente discriminantes, permitindo que os modelos fossem capaz de identificar quais recortes continham a planta invasora e quais travam-se de amostras de pastagem.

Os resultados mostraram que todas as estratégias avaliadas foram capazes de alcançar taxas de acerto superiores a 90%. O modelo que apresentou a maior acurácia para o conjunto em que havia distribuição igual entre amostras negativas e positivas foi a combinação pela regra do produto, com 95,55% de acertos. Para o segundo conjunto de dados (Conjunto 2.0), no entanto, o melhor modelo foi o KNN, com acurácia de 95,92%.

Ao compararmos o desempenho dos modelos nos dois conjuntos de dados em termos de acurácia, percebe-se pouca variação. No entanto, como esperado, os índices de sensibilidade apresentaram variações maiores. Ao aumentarmos o número de exemplares da classe negativa, para tentar representar um cenário mais realista, percebemos que os modelos RF e SVM apresentaram uma diminuição em sua capacidade de detectar a classe positiva em 4,48 e 5,61 pontos percentuais, respectivamente. Os dois modelos que tiveram mais sucesso em detectar a presença da *Rumex* foram o KNN e o MLP, com taxas de sensibilidade de 0,9003 e 0,8886, respectivamente.

Portanto é perceptível que escolher uma base de dados, ou um modelo de aprendizado com base apenas na acurácia, pode não oferecer a melhor solução, pois alguns problemas podem demandar uma taxa mínima de acertos para classes específicas. No caso da identificação de plantas invasoras, por exemplo, identificar uma erva daninha onde não ela não se faz presente (erro de comissão), traz menos danos do que um erro de omissão. Isso ocorre porque aplicar herbicida na grama, por engano, é menos prejudicial do que deixar de aplicá-lo nas plantas daninhas, que podem causar maiores prejuízos, principalmente a longo prazo.

Os próximos passos da pesquisa envolvem a adoção de novos conjuntos de dados contendo imagens de plantas invasoras, com características distintas de aquisição, como as apresentadas no trabalho de [Güldenring et al. \(2023\)](#) em que os autores construíram um conjunto de imagens de pastagens com a presença de *Rumex obtusifolius* e *Rumex crispus*. Afora a avaliação individualizada das imagens espera-se, futuramente, fazer a adoção de conjuntos de imagens geograficamente referenciadas, de forma que seja possível a construção de mosaicos que identifiquem exatamente onde estão as plantas invasoras para o combate otimizado.

Além de conjuntos de dados variados, é pertinente a avaliação de estratégias de aprendizado profundo, uma vez

que trabalhos presentes na literatura indicaram a pertinência de tais técnicas, como observado em (Wu et al., 2021). Nesse contexto, pretende-se adotar algoritmos como as redes neurais convolucionais clássicas (*Convolutional Neural Networks*) e variações de sua implementação, como a ResNet (*Residual Network*), SegNet, VGGNet (*Visual Geometry Group*), GoogleNet, FCN (*Fully Convolutional Network*) e YOLO (*You Only Look Once*).

Referências

- Aha, W., Kibler, D. and Albert, M. (1991). Instance-based learning algorithms, *Machine Learning* 6: 37–66. <https://doi.org/10.1023/A:1022689900470>.
- Aragão, A. and Contini, E. (2022). O agro no brasil e no mundo: uma síntese do período de 2000 a 2020, *Technical report*, Empresa Brasileira de Pesquisa Agropecuária (Embrapa), Brasília, Brasil. Disponível em <https://www.embrapa.br/documents/10180/62618376/0+AGRO+NO+BRASIL+E+NO+MUNDO.pdf/41e20155-5cd9-f4ad-7119-945e147396cb>.
- Araújo, V., Britto, A. S., Brun, A. L., Koerich, A. L. and Palate, R. (2017). Multiple classifier system for plant leaf recognition, *2017 IEEE International Conference on Systems, Man, and Cybernetics (SMC)*, pp. 1880–1885. <https://doi.org/10.1109/SMC.2017.8122891>.
- Breiman, L. (2001). Random forests, *Machine Learning* 45: 5–32. <https://doi.org/10.1023/A:1010933404324>.
- Castro, L. N. and Ferrari, D. G. (2016). *Introdução à Mineração de dados: Conceitos básicos, algoritmos e aplicações*, 1 edn, Saraiva, São Paulo, SP, Brasil.
- Espejo-Garcia, B., Malounas, I., Vali, E. and Fountas, S. (2021). Testing the suitability of automated machine learning for weeds identification, *AI* 2(1): 34–47. <https://doi.org/10.3390/ai2010004>.
- Esposito, M., Crimaldi, M., Cirillo, V., Sarghini, F. and Maggio, A. (2021). Drone and sensor technology for sustainable weed management: a review, *Chemical and Biological Technologies in Agriculture* 8(18): 1–11. <https://doi.org/10.1186/s40538-021-00217-8>.
- Faceli, K., Lorena, A. C., Gama, J. and de Carvalho, A. C. P. L. F. (2011). *Inteligência Artificial: uma abordagem de aprendizado de máquina*, 1 edn, Editora LTC, Rio de Janeiro, RJ, Brasil.
- Gunes, V., Menard, M., Loonis, P. and Petitrenaud, S. (2003). Combination, cooperation and selection of classifiers: A state of the art., *International Journal of Pattern Recognition and Artificial Intelligence* 17: 1303–1324. <https://doi.org/10.1142/S0218001403002897>.
- Güldenring, R., van Evert, F. K. and Nalpantidis, L. (2023). Rumexweeds: A grassland dataset for agricultural robotics, *Journal of Field Robotics* 40(6): 1639–1656. <https://doi.org/10.1002/rob.22196>.
- Güttler, G., Milani, V., Tarouco, C. P., Nohatto, M. A. and Rosa, E. d. F. d. F. (2018). Estratégias de controle físico-mecânico sobre rumex obtusifolius, *Revista Ciência Agrícola* 16(2): 41–45. <https://doi.org/10.28998/rca.v16i2.4117>.
- IBGE (2022). Produção agropecuária. Disponível em <https://www.ibge.gov.br/explica/producao-agropecuaria/>.
- Jochem, W., Visentin, R. P., Nogatz, B., Costa, G. D., Schmitt, J., de Oliveira Neto, A. M. and Guerra, N. (2020). The perceptron: a probabilistic model for information storage and organization in the brain, *Colloquium Agrariae* 16(5): 127–134. Disponível em <https://revistas.unoeste.br/index.php/ca/article/view/3241>.
- Li, Y., Al-Sarayreh, M., Irie, K., Hackell, D., Bourdot, G., Reis, M. M. and Ghamkhar, K. (2021). Identification of weeds based on hyperspectral imaging and machine learning, *Frontiers in Plant Science* 11: 1–13. <https://doi.org/10.3389/fpls.2020.611622>.
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V., Vanderplas, J., Passos, A., Cournapeau, D., Brucher, M., Perrot, M. and Duchesnay, E. (2011). Scikit-learn: Machine learning in Python, *Journal of Machine Learning Research* 12: 2825–2830. <https://dl.acm.org/doi/10.5555/1953048.2078195>.
- Pereira, F. A. R., Verzignassi, J. R., Arias, E. R. A., Carvalho, F. T. and Paula e Silva, A. (2011). Documentos 185: Controle de plantas daninhas em pastagens, *Technical report*, Empresa Brasileira de Pesquisa Agropecuária (Embrapa), Brasília, Brasil. Disponível em <https://www.infoteca.cnptia.embrapa.br/infoteca/bitstream/doc/920044/1/D0C1851.pdf>.
- Platt, J. (1998). Sequential minimal optimization: A fast algorithm for training support vector machines, *Advances in Kernel Methods-Support Vector Learning* 208. Disponível em <https://www.researchgate.net/publication/2624239>.
- Portal do Agronegócio (2022). Manejo eficiente da pastagem ajuda na produção de carne de qualidade. Disponível em <https://tinyurl.com/2wh8neaj/>.
- Sapkota, B. B., Popescu, S., Rajan, N., Leon, R. G., Reberg-Horton, C. and Mirsky, S. (2022). Use of synthetic images for training a deep learning model for weed detection and biomass estimation in cotton, *Scientific Reports* 12: 1–18. <https://doi.org/10.1038/s41598-022-23399-z>.
- Smith, L. N., Byrne, A., Hansen, M. F., Zhang, W. and Smith, M. L. (2019). Weed classification in grasslands using convolutional neural networks, in M. E. Zelinski, T. M. Taha, J. Howe, A. A. S. Awwal and K. M. Iftekharrudin (eds), *Applications of Machine Learning*, Vol. 11139, International Society for Optics and Photonics, SPIE, San Diego, California, United States, pp. 1–11. <http://dx.doi.org/10.1117/12.2530092>.
- Terada, Y. (2022). rumex_detection_no10 dataset. Disponível em https://universe.roboflow.com/yoshiki-terada/rumex_detection_no10.

- Victoria Filho, R., Ladeira Neto, A., Pelissari, A., Reis, F. C. and Daltro, F. P. (2014). Manejo sustentável de plantas daninhas em pastagens, *Manejo de Plantas Daninhas nas Culturas Agrícolas*, RiMa, São Carlos, São Paulo, pp. 179–207.
- Vilar, D. (2021). Desafios no manejo sustentável de plantas daninhas em pastagens. Disponível em <https://agronline.com.br/portal/artigo/desafios-no-manejo-sustentavel-de-plantas-daninhas-em-pastagens/>.
- Wu, Z., Chen, Y., Zhao, B., Kang, X. and Ding, Y. (2021). Review of weed detection methods based on computer vision, *Sensors* 21(11): 1–23. <https://doi.org/10.3390/s21113647>.
- Zaller, J. (2006). Allelopathic effects of rumex obtusifolius leaf extracts against native grassland species, *Journal of Plant Diseases and Protection* 20: 463–470. Disponível em <https://tinyurl.com/2tcury5x>.
- Zhang, W., Hansen, M. F., Volonakis, T. N., Smith, M., Smith, L., Wilson, J., Ralston, G., Broadbent, L. and Wright, G. (2018). Broad-leaf weed detection in pasture, *2018 IEEE 3rd International Conference on Image, Vision and Computing (ICIVC)*, Chongqing, China, pp. 101–105. <https://doi.org/10.1109/ICIVC.2018.8492831>.