



UNIVERSIDADE FEDERAL DE PERNAMBUCO
CENTRO DE INFORMÁTICA
PROGRAMA DE PÓS-GRADUAÇÃO EM CIÊNCIA DA COMPUTAÇÃO

THIAGO MOURA DA ROCHA BASTOS

Avaliação de Filmes Metalizados Por Algoritmos de Aprendizagem de Máquina Através de
Dados Operacionais de Processo Industrial e de Qualidade.

Recife

2021

THIAGO MOURA DA ROCHA BASTOS

Avaliação de Filmes Metalizados Por Algoritmos de Aprendizagem de Máquina Através de Dados Operacionais de Processo Industrial e de Qualidade.

Trabalho apresentado ao Programa de Pós-graduação em Ciência da Computação do Centro de Informática da Universidade Federal de Pernambuco, como requisito parcial para obtenção do título de Mestre em Ciência da Computação.

Área de Concentração: Inteligência Computacional

Orientador: Cleber Zanchettin

Coorientador: Luiz Stragevitch

Recife

2021

Catálogo na fonte
Bibliotecária Monick Raquel Silvestre da S. Portes, CRB4-1217

B327a Bastos, Thiago Moura da Rocha
Avaliação de filmes metalizados por algoritmos de aprendizagem de máquina através de dados operacionais de processo industrial e de qualidade / Thiago Moura da Rocha Bastos. – 2021.
86 f.: il., fig., tab.

Orientador: Cleber Zanchettin.
Dissertação (Mestrado) – Universidade Federal de Pernambuco. CIn, Ciência da Computação, Recife, 2021.
Inclui referências.

1. Inteligência computacional. 2. Aprendizagem de máquina. I. Zanchettin, Cleber (orientador). II. Título.

006.31 CDD (23. ed.) UFPE - CCEN 2021 – 137

Thiago Moura da Rocha Bastos

“Avaliação de Filmes Metalizados por Algoritmos de Aprendizagem de Máquina Através de Dados Operacionais de Processo Industrial e de Qualidade”

Dissertação de Mestrado apresentada ao Programa de Pós-Graduação em Ciência da Computação da Universidade Federal de Pernambuco, como requisito parcial para a obtenção do título de Mestre em Ciência da Computação. Área de Concentração: Inteligência Computacional

Aprovado em: 30/06/2021.

BANCA EXAMINADORA

Prof. Dr. Leandro Maciel Almeida
Centro de Informática / UFPE

Prof. Dr. Frederico Duarte de Menezes
Departamento de Sistemas, Processos e Controles Industriais / IFPE

Prof. Dr. Cleber Zanchettin
Centro de Informática/ UFPE
(Orientador)

Dedico esse texto a quem tem sede e fome de justiça, a quem sonha com oportunidades para todos e aos que se inquietam diante da falta dela, e a quem valoriza o ensino público, de qualidade e universal. Pois, só quem sofre com a desigualdade é capaz de saber o quanto ela nos limita, constrange e incapacita.

AGRADECIMENTOS

Agradeço a todas as pessoas que de alguma forma acreditaram, motivaram e contribuíram para o desenvolvimento desse trabalho. Em especial, a minha família, por todo suporte e incentivo. Ao meu orientador, Prof. Cleber Zanchettin, que desde sempre se mostrou disponível para apoiar, colaborar e melhorar nossas ideias para o projeto. Gostaria de agradecer também ao Centro de informática (Cin) e ao Laboratório de Combustíveis (LAC) pelo suporte acadêmico, disponibilidade de espaço e pelas parcerias de projetos.

É evidente que a engenharia química está em uma encruzilhada importante. Nossa disciplina está passando por uma transição sem precedentes que apresenta desafios e oportunidades significativos em modelagem e tomada de decisão automatizada. (VENKATASUBRAMANIAN, 2019, p.1, tradução nossa).

RESUMO

Os processos industriais de manufatura são parte essencial da revolução tecnológica vivida na atualidade. Dentre esses, a produção de filmes metalizados por deposição contribuem para embalagens, utilizadas principalmente na conservação de alimentos. Com crescimento global médio de 10-15% por ano desse mercado, modelos tradicionais de monitoramento e controle tem representado gargalos para a entrega desses produtos além de elevadas taxas de produtos reprovados, e o uso de novas tecnologias digitais disruptivas como a inteligência artificial pode ser uma alternativa para superar essas dificuldades. Assim, esse trabalho objetiva utilizar sistemas de aprendizado de máquina para interpretação e predição de variáveis de processo e qualidade presentes na produção de filmes metalizados por deposição à vácuo e redução do tempo atual de entrega dos produtos finalizados. Comparando diferentes classificadores de aprendizado de máquina associados a condições diversas de préprocessamento de dados e hiper-parâmetros para a predição de qualidade do produto, o modelo *Random Forest* apresentou o maior desempenho com 85,4% de acurácia. Foram utilizadas diferentes técnicas de visualização para interpretar as previsões e observar o desempenho dos modelos aplicados. Por outro lado, através da segmentação semântica dos perfis de densidade óptica dos produtos, foi possível a identificação de falhas, e o monitoramento da qualidade final dos filmes produzidos por um modelo de rede neural com 86,67% de acurácia. Além disso, a aplicação das visualizações auxiliam no entendimento e validação dos produtos obtidos no processo de metalização e associados a diferentes condições operacionais sobre os produtos manufaturados. Este estudo de caso demonstra o potencial uso de modelos de aprendizado de máquina para suporte a analistas e operadores na interpretação de variáveis operacionais, oferecendo informações relevantes para monitorar e manter o processo de metalização de filme por deposição a vácuo e servir como base para análises de desempenho e robustez no futuro implementação industrial.

Palavras-chaves: inteligência artificial; aprendizagem de máquina; visualização; monitoramento de processo; predição de qualidade.

ABSTRACT

Industrial manufacturing processes are an essential part of the technological revolution experienced today. Among these, the production of metallized films by deposition contributes to packaging, mainly used in food preservation. With an average global growth of 10-15 % per year in this market, traditional monitoring and control models have represented bottlenecks for the delivery of these products, in addition to high rates of products that need to be reprocessed. The use of new digital breakthrough technologies as artificial intelligence can be an alternative to overcome these difficulties. Thus, this work aims to use machine learning systems to interpret and predict process and quality variables present in the production of metallized films by vacuum deposition and reduce the actual delivery time of finished products. Comparing different machine learning classifiers associated with different data preprocessing conditions and hyper-parameters to predict product quality, the *Random Forest* model showed the highest performance with 85.4% accuracy. Were used visualization techniques to interpret these predictions and observe the model's performance. On the other hand, through the semantic segmentation of the optical density profiles of the products, it was possible to identify flaws and monitor the final quality of the films produced by a neural network model with 86.67% accuracy. In addition, the application of visualizations helps in the understanding and validation of the products obtained in the metallization process and associated with different operational conditions on the manufactured products. This case study demonstrates the potential use of machine learning models to support analysts and operators in interpreting operational variables, offering relevant information to monitor and maintain the vacuum metallization process and serve as a basis for analysis of performance and robustness in future industrial deployment.

Keywords: artificial intelligence; machine learning; visualization; process monitoring; quality prediction.

LISTA DE FIGURAS

- Figura 1 – Imagem do produto final do processo industrial de metalização a vácuo. O filme metalizado pode apresentar dimensões variadas, com até 4,45 m de largura e até 85 km de extensão. 26
- Figura 2 – Ilustração da zona de deposição presente no processo de metalização a vácuo. Destacando as barcas cerâmicas dispostas transversalmente em relação ao comprimento do filme, e a alimentação de fios metálicos por bobinas de passo para liquefação e vaporização do metal pelas barcas aquecidas. Por fim, o metal vaporizado é depositado sobre a superfície do filme. 26
- Figura 3 – Efeito da energia superficial sobre a qualidade do recobrimento. A depender das condições da superfície do filme para metalização, pode ocorrer a formação de filmes com altos valores de densidade óptica, porém com recobrimento pobre. 28
- Figura 4 – Ocorrência de áreas de diferentes aspectos causadas pelo desgaste em uma barca de evaporação metálica. Essa desuniformidade na superfície cerâmica pode provocar o surgimento de falhas no recobrimento do filme na sua respectiva posição, levando a formação de produtos de baixa qualidade. . . 29
- Figura 5 – Dependência do perfil da nuvem metálica de deposição de acordo com a temperatura da barca. Esse feito demonstra a importância direta da estabilidade da temperatura presente nas barcas cerâmicas, além da necessidade que todo o conjunto de barcas apresente performance semelhante de estabilidade de temperatura, de forma a manter uniformidade no recobrimento metálico do filme. 30
- Figura 6 – Imagem da metalizadora TopMet2450 utilizada para processo industrial de recobrimento metálico em filme polimérico. 32

Figura 7 – Ilustração do processo de metalização à vácuo. O processo inicia com a alimentação de um filme polimérico (substrato) a ser desenrolado, guiado até a zona de deposição e recoberto pelo metal vaporizado em barcas cerâmicas. Esse metal é alimentado em forma de fios por bobinas de passo, para serem liquefeitos, vaporizados e depositados sobre a superfície do filme. Após o recobrimento, através de sensores são obtidas informações de qualidade do recobrimento. E por fim o filme recoberto é enrolado novamente.	34
Figura 8 – Subdivisão dos modelos de aprendizagem de máquina de acordo com o seu tipo de aprendizado e suas principais aplicações.	36
Figura 9 – Tipos de predição para modelos supervisionados de aprendizagem de máquina. A predição de classe tem como objetivo a classificação de novos objetos em um espaço com pelo menos duas classes representadas. Já a regressão é feita através da modelagem de exemplos pertencentes a mesma classe para obtenção de valores de variáveis para novos objetos.	38
Figura 10 – Ilustração de um perceptron, estrutura básica do algoritmo de aprendizado profundo baseado em redes neurais, o perceptron multicamadas (<i>Multi-Layer Perceptron</i>) (MLP). As camadas de entrada são as variáveis X_i e o Y são os rótulos preditos pelo modelo. Já os parâmetros W_i e a função de ativação Σ são, respectivamente, os parâmetros e o hiper-parâmetros do modelo.	39
Figura 11 – Função de ativação tangente hiperbólica, $\tanh$, utilizada em modelos de ANN para modelagem de problemas não-lineares. A não-linearidade dos modelos é obtida através da transformação dos valores y fornecidos pelos perceptrons em valores de $\tanh y$	40
Figura 12 – Função de ativação unidade linear retificada (<i>Rectified Linear Unit</i>) (ReLU) utilizada em modelos de ANN para modelagem de problemas não-lineares. Para valores negativos a função retorna zero, e valores positivos a função tem comportamento linear com inclinação igual a um.	40

- Figura 13 – Diferenciação do processo de modelagem por modelos estatísticos e por modelos de aprendizado de máquina (*Machine Learning*) (ML). Os modelos estatísticos utilizam de funções matemáticas para representação das variáveis, cujo o objetivo é a redução do erro entre a representação dos exemplos e os próprios exemplos. Já os modelos de ML tem como função objetivo a minimização do erro entre valores de predição e os valores verdadeiros atribuídos aos objetos de um espaço, utilizando para isso o ajuste de parâmetros e hiper-parâmetros construídos e fornecidos, respectivamente. 42
- Figura 14 – Modelos estatísticos utilizam, em geral, uma subdivisão dos dados para análise de multicolinearidade e outra parte para análise de ajuste de funções paramétricas. Por outro lado, os modelos de aprendizado de máquina são construídos com a subdivisão de treinamento, ajustados com a subdivisão de validação e verificados com a subdivisão de teste, o que fornece mais robustez para generalização do modelos com novos exemplos. 43
- Figura 15 – Construção de sistemas de modelagem em múltiplas escalas e modelos de aprendizado de máquina (ML) para acelerar criação de sistemas. A combinação desses dois fatores facilita a descoberta de modelos interpretáveis a partir de dados heterogêneos e irregulares, guiando a aquisição e interpretação física-química de novas informações. 44
- Figura 16 – Espaço de configurações do Auto-Weka. Primeiramente é selecionado o método de classificação, a partir dos 27 modelos básicos, *base models*, de ML disponíveis (triângulos) ou por modelos com configurações específicas a serem definidas (elipses). Já a segunda etapa consiste na seleção das etapas de seleção e pre-processamento das variáveis. 46
- Figura 17 – Representação do espaço de busca do modelo automático Hyperopt-sklearn, que consiste na conjunção de aplicações disponíveis no sklearn de pre-processamento e de classificadores. Com 6 opções de cada, o hyperopt constrói suas opções de modelos a partir de amostragens aleatórias. Os blocos em verde mostram um exemplo de modelo construído, onde esse modelo (PCA e KNN) representam em conjunto o uso de 7 módulos do Hyperopt. . . . 47

Figura 18 – Ilustração das etapas de construção de um sistema de ML através do modelo automático TPOT. Esse modelo automático busca a otimização das etapas de processamento de dados, de variáveis, da escolha dos modelos de ML e da otimização de hiper-parâmetros através de modelos baseados em estruturas de árvores.	48
Figura 19 – Espaço de configurações do modelo de implementação automática Auto-Sklearn. Através da análise de performance e da otimização bayesiana de hiper-parâmetros, são construídos os sistemas de predição, dados pela combinação das etapas de processamento de dados, variáveis e pela escolha do modelo de ML. Em azul são os hiper-parâmetros principais, já em cinza são os hiper-parâmetros secundários escolhidos para exemplo de possível configuração para um sistema de ML.	49
Figura 20 – Exemplo de segmentação semântica em quatro classes com divisão de imagem em super-píxeis e agrupamento por classes através de um classificador de árvores aleatórias, utilizando de imagens de infravermelho de ovos de <i>Drosophila</i> obtidas em microscópio eletrônico. A primeira coluna de imagens mostra as imagens originais, a segunda a segmentação em classes representadas por diferentes cores e a terceira coluna demonstra a construção supervisionada dos seguimentos e a localização de objetos presentes na imagem.	53
Figura 21 – Matriz de correlação com variáveis de processo com correlação maior que 97,5%. A presença de variáveis com correlações altas é dada, principalmente, pelos valores de densidade ópticas obtidos nas diferentes posições dos sensores. Isso pode ser entendido, uma vez que esses valores de recobrimento variam em torno de um mesmo valor pré-determinado, fazendo com que os valores de OD acabem apresentando variações coordenadas, mesmo com os sistemas locais de recobrimento agindo de forma independente. . .	57
Figura 22 – Distribuição dos valores de densidade óptica dos objetos. Sendo cada coluna referente a posição do sistema de deposição onde há o recobrimento metálico e também as medidas de OD ao longo do comprimento do filme. .	59

Figura 23 – Pipeline Profiler aplicado para visualização de sistemas de predição de classe dos filmes obtidos pelo Auto-sklearn aplicado a variáveis do processo de metalização a vácuo. A visualização é composta por: 1) Sistemas do Auto-sklearn formados pela combinação dos modelos primitivos de processamento de dados, classificadores e balanceamento, descritos pelas colunas e símbolos e ordenados de acordo com a sua performance. 2) Descrição dos hiper-parâmetros utilizados para o sistema selecionado. 3) Valores de performance de acordo com os sistemas construídos. Sendo também possível observar os modelos ordenados pelo tempo de predição. 4) Contribuição de cada etapa para a performance do sistema.	62
Figura 24 – Representação dos dez sistemas com maior desempenho fornecidos pelo Auto-sklearn. As configurações dos sistemas de maiores desempenhos demonstram o uso de diferentes combinações de modelo de ML, pré-processamento de variáveis e de pré-processamento de dados, com os respectivos valores de desempenho.	63
Figura 25 – Análise dos dez sistemas com maiores acurácias fornecidos pelo Auto-sklearn por diferentes métricas de performance. Foram escolhidas as métricas comumente utilizadas para análise de modelos de ML.	64
Figura 26 – Distribuição de probabilidade de predição dos classificadores RF e LDA. A comparação entre as distribuições mostra que existe uma maior quantidade de erros de falso positivo ($p > 0.5$) para o modelo LDA.	65
Figura 27 – Importância das variáveis para os classificadores RF1 e LDA, de acordo com os valores SHAP médios de impacto. O que permite a observância entre a importância real das variáveis do processo industrial e como elas afetam no processo de predição de qualidade.	66
Figura 28 – Dispersão dos valores das variáveis do processo que compõem os objetos, e os seus respectivos impactos sobre o modelo RF para predição de filmes aprovados. Demonstrando a correlação existente entre a discriminação dos valores das variáveis e a sua significância para a construção do modelo. . .	67

- Figura 29 – Modelo de rede neural sequencial do Keras adaptada para classificação dos filmes metalizados utilizando informações dos perfis de OD. O modelo é composto por duas camadas convolucionais alternadas de duas camadas de conjugação 2x2. Em seguida é aplicada uma regularização de 0.5 para evitar sobreajuste do modelo, e em seguida a linearização e determinação da classe dos objetos pela função de ativação softmax. 68
- Figura 30 – a) Perfil de OD do filme metalizado. Valores SHAP de impacto sobre a predição das classes b) Aprovado e c) Reprovado. Os valores positivos de SHAP, marcados em vermelho, favorecem a classe observada, já os valores em azul prejudicam a classificação observada. O que permite a visualização da distribuição desses valores de SHAP de acordo com a qualidade do recobrimento. 69
- Figura 31 – Segmentação semântica em quatro classes dos perfis de OD dos filmes para identificação de falhas. As diferentes cores representam as classes obtidas por agrupamento, e a área em azul marinho representa área onde não há filme. De acordo com o formato e a distribuição das classes, é possível inferir que a classe 1, azul claro, é uma provável região de bom recobrimento, a classe 2, amarelo, é uma provável área de recobrimento intermediário, e por fim a classe 3, em vermelho, como representante de regiões de recobrimento pobre. a) Perfil segmentado de um objeto aprovado. b) Perfil segmentado de um objeto reprovado. 71
- Figura 32 – Demonstração dos valores de contagem de acordo com as classes utilizadas na segmentação semântica dos perfis de OD dos filmes. As barras pretas verticais demonstram os valores de contagem para os diferentes exemplos cuja classificação é descrita de acordo com a cor da barra, verde para filmes aprovados e amarelo para filmes reprovados. 72
- Figura 33 – Análise de sensibilidade e espacialidade de informações a partir das saídas das três primeiras camadas do modelo de rede neural treinados com imagens dos perfis de OD originais (a esquerda) e com os perfis após segmentação semântica (a direita). a) Mapeamento de variáveis utilizando objetos com perfis de OD originais. b) Mapeamento de variáveis utilizando objetos com perfis de OD após segmentação semântica. 73

- Figura 34 – Utilização de plotes paralelos com variáveis de processo mais relevantes para a predição pelo modelo RF. O que permite a observação do histórico de correlação das variáveis para objetos aprovados, linhas amarelas, e para objetos reprovados, linhas azuis. Mostrando a existência de correlações diretas e correlações inversas entre variáveis vizinhas, o que pode auxiliar operadores na escolha de parâmetros operacionais que favoreçam maiores taxas de produtos aprovados. 75
- Figura 35 – Utilização de plotes paralelos com valores de variáveis de processo mais relevantes restritos a certas condições operacionais. Demonstrando que há influência da combinação de valores das variáveis de processo sobre a taxa de produtos aprovados. E que a certas condições operacionais devem ser evitadas, já que apresentam uma elevada taxa de produtos reprovados. . . 75
- Figura 36 – Resumo de predição de classe do modelo de ML para um exemplo aprovado. Essa visualização demonstra os impactos das variáveis sobre a classificação do produto, e a descrição das condições operacionais durante a sua manufatura. Tendo como objetivo, permitir operadores e analistas a observação e interpretação dos produtos e da sua classificação em tempo de produção. 76

LISTA DE TABELAS

- Tabela 1 – Descrição do banco de dados formado a partir das variáveis do processo industrial de metalização de filme por deposição. As colunas apresentam os valores médios e os valores presentes nos quantis. O que permite a observação das condições operacionais utilizadas na obtenção dos filmes metalizados. A variável situação é o rótulo correspondente a classificação final dos produtos, sendo o valor um para produtos aprovados e zero para produtos reprovados. 58
- Tabela 2 – Descrição da base de dados formada com valores de densidade óptica (*Optical Density*) (OD). Os valores foram obtidos a partir de 26 sensores ópticos distribuídos transversalmente ao comprimento do filme, com registro de informação a cada dois minutos e variação de zero a cinco, o que representa o menor e o maior valor de recobrimento, respectivamente. Associado aos valores de OD também são obtidos os respectivos valores de comprimento das medições. 59

LISTA DE ABREVIATURAS E SIGLAS

AdaBoost	classificador de impulso adaptativo (<i>Adaptive Boosting classifier</i>)
AI	inteligência artificial (<i>Artificial Intelligence</i>)
ANN	redes neurais artificiais (<i>Artificial Neural Networks</i>)
AO	otimização de arquitetura (<i>Architecture Optimization</i>)
AUC	área debaixo da curva (<i>Area Under the Curve</i>)
AutoML	aprendizado de máquina automático (<i>Automatic Machine Learning</i>)
BO	otimização Bayesiana (<i>Bayesian Optimization</i>)
CASH	seleção combinada de algoritmos e de hiper-parâmetros (<i>Combined Algorithm Selection and Hyperparameter Optimization</i>)
Cin	Centro de informática
DL	aprendizado profundo (<i>Deep Learning</i>)
ExTree	classificador de Árvores Extras (<i>Extra Trees classifier</i>)
GMM	modelo de mistura de gaussianas (<i>Gaussian Mixture Model</i>)
HPO	otimização de hiper-parâmetros (<i>Hyperparameters Optimization</i>)
IoT	internet das coisas (<i>Internet of Things</i>)
KNN	k-vizinhos mais próximos (<i>K-Nearest Neighbors</i>)
LAC	Laboratório de Combustíveis
LDA	análise por discriminante linear (<i>Linear Discriminant Analysis</i>)
ML	aprendizado de máquina (<i>Machine Learning</i>)
MLP	perceptron multicamadas (<i>Multi-Layer Perceptron</i>)
MPC	modelo de Controle Preditivo (<i>Model Predictive Control</i>)
OD	densidade óptica (<i>Optical Density</i>)
PCA	análise por componentes principais (<i>Principal Component Analysis</i>)
QDA	análise por discriminante quadrático (<i>Quadratic Discriminant Analysis</i>)
ReLU	unidade linear retificada (<i>Rectified Linear Unit</i>)

RENAS	pesquisa por arquitetura neural evolutiva reforçada (<i>Reinforced Evolutionary Neural Architecture Search</i>)
RF	florestas aleatórias (<i>Random Forest</i>)
SHAP	explicação aditiva de Shapley (<i>SHapley Additive Explanations</i>)
SS	segmentação semântica (<i>Semantic Segmentation</i>)
SVC	classificador de vetores de suporte (<i>Support Vector Classifier</i>)
SVM	máquina de Vetores de Suporte (<i>Support Vector Machine</i>)
TPOT	ferramenta de otimização com estrutura baseada em árvores (<i>Tree-Based Pipeline Optimization Tool</i>)
TSNE	incorporação estocástica de vizinho distribuído em T (<i>T-distributed Stochastic Neighbor Embedding</i>)
UMAP	aproximação e projeção uniforme de múltiplas dobras (<i>Uniform Manifold Approximation and Projection</i>)

SUMÁRIO

1	INTRODUÇÃO	21
1.1	OBJETIVOS	23
1.1.1	Objetivo geral	23
1.1.1.1	<i>Objetivos específicos</i>	23
1.2	CONTRIBUIÇÕES	23
1.3	ORGANIZAÇÃO DO TRABALHO	24
2	FUNDAMENTAÇÃO TEÓRICA	25
2.1	PROCESSO INDUSTRIAL DE RECOBRIMENTO	25
2.1.1	Recobrimento de filmes poliméricos	25
2.1.2	Desafios do processo de recobrimento metálico de filmes por deposição	27
2.1.3	Variáveis de processo e de qualidade	28
2.1.4	Processo industrial de deposição metálica em filme	31
2.2	INTELIGÊNCIA ARTIFICIAL	34
2.2.1	Algoritmos de aprendizagem de máquina	35
2.2.1.1	<i>Aprendizagem supervisionada</i>	38
2.2.1.2	<i>Modelos estatísticos x modelos de ML</i>	41
2.2.1.3	<i>Algoritmos de ML e os processos industriais</i>	43
2.3	IMPLEMENTAÇÃO AUTOMÁTICA DE MODELOS DE AM	45
2.3.1	Auto-Sklearn	48
2.4	VISUALIZAÇÃO DE VARIÁVEIS E MODELOS DE ML	50
2.4.1	Métodos de interpretação de modelos de ML e variáveis	50
2.4.2	Métodos de visualização de objetos	52
2.4.2.1	<i>Segmentação semântica</i>	52
3	METODOLOGIA	55
3.1	COLETA DE DADOS	55
3.2	PREPROCESSAMENTO DOS DADOS	55
3.2.1	Seleção de variáveis	56
3.3	TREINAMENTO DOS ALGORITMOS	60
3.4	VISUALIZAÇÃO DOS RESULTADOS	60

4	RESULTADOS E DISCUSSÃO	62
4.1	AUTOML	62
4.2	ANÁLISE DOS FILMES POR VISUALIZAÇÃO DO PERFIL DE DO	68
4.3	INTERPRETAÇÃO DE VARIÁVEIS DE PROCESSO E DO PRODUTO . .	74
5	CONCLUSÃO	78
	REFERÊNCIAS	79

1 INTRODUÇÃO

A demanda mundial por crescimento econômico e aumento de produtividade através de um desenvolvimento sustentável, exige de governos e empresas avanços tecnológicos, inovações em suas dinâmicas de processos e uma lógica produtiva sob demanda para redução de desperdícios (OVERCASH, 2019). Exigências essas que se tornam ainda mais desafiadoras em um contexto global de informações superpostas, com características singulares de grandes volumes, veracidade, variedade, valor, visibilidade e velocidade, Vs que representam o termo *Big Data* presente no contexto das indústrias 4.0 (LEE; SHIN; REALFF, 2018). Os processos industriais de manufatura compõem parte importante dessa nova dinâmica global, e as indústrias que aplicam novas tecnologias para essas demandas são conhecidas como manufaturas avançadas.

A construção de manufaturas avançadas através da adoção de sistemas inteligentes na área industrial, e mais especificamente, no desenvolvimento de embalagens e filmes recobertos inteligentes, pode ser útil para estender o tempo de prateleira de alimentos, melhorar a qualidade e segurança das embalagens, e prover informações importantes sobre o produto (BIJI et al., 2015). São encontradas diversas aplicações de filmes e produtos com recobrimento metálico, seja com objetivos decorativos como as embalagens brilhosas ou decoradas através de impressões sobre o metal, com objetivos de barreira para luz, vapor e gases, com objetivos funcionais como as embalagens semicondutoras utilizadas no aquecimento de alimentos, ou mesmo com objetivo de segurança, como o caso de etiquetas holográficas (BISHOP; III, 2016). Além disso, no caso de processos industriais de recobrimentos metálicos, Joshi (2020) aponta para o uso de inteligência artificial como substituto de sistemas tradicionais de controle, e mais precisamente, o uso de redes neurais para análise e controle de variáveis envolvidas no processo.

A análise atual de filmes metalizados pelo processo industrial de recobrimento a vácuo apresenta limitações, sendo feita através de pequenos recortes do produto e de forma laboratorial, com equipamentos custosos e com análises que demandam grande períodos de espera. O que torna a dinâmica de produção lenta, já que os produtos não podem ser liberados para venda antes da análise de qualidade. Além disso, as condições operacionais são poucos utilizadas no monitoramento dos produtos, uma vez que não existem sistemas que demonstrem o impacto das variáveis sobre o processo, o que pode levar a formação de produtos de baixa qualidade e a perda de materiais por elevadas taxas de produtos que necessitam de repro-

cessamento. Essas dificuldades corroboram com a necessidade da construção de um sistema de análise sistemática que contribua na observação de qualidade dos produtos para além da análise laboratorial, e de forma paralela, permita a interpretação das variáveis que impactam na qualidade dos produtos.

Para esses tipos de problemas, os processos de manufaturas avançadas já tem adotado tecnologias baseadas em informações em tempo real, manutenções preditivas, decisões baseadas em dados e força de trabalho altamente qualificada (DAVISA et al., 2012). Essas aplicações só são possíveis através do uso massivo de análise de dados operacionais obtidos por sensores, modelagens e simulações de processos, que coletam dados e através de memórias e processadores de alta performance integram sistemas para obtenção de informações através de aplicações de internet das coisas (*Internet of Things*) (IoT) (WORTMANN; FLÜCHTER, 2015).

Em complemento aos métodos de ML, recentes técnicas de visualização e processamento de dados como Theano (BERGSTRA et al., 2010), Pylearn2 (GOODFELLOW; WARDE-FARLEY, 2013), Caffe (JIA et al., 2014) e Torch (COLLOBERT; KAVUKCUOGLU, 2011), buscam agregar à eficiência preditiva (LECUN; BENGIO; HINTON, 2015), uma capacidade interpretativa a esses modelos (Ming; Qu; Bertini, 2019). Auxiliando atividades industriais como melhoria de processo, identificação de falhas e padrões (AMERSHI et al., 2015; KAHNG et al., 2017; Wongsuphasawat et al., 2018).

Além disso, aplicações de ML em visão computacional tem demonstrado capacidade de localizar e identificar objetos em imagens, assim como contar e extrair padrões para criação de classes e compreensão de padrões, seja por interpretação de pesos presentes nos modelos de redes neurais ou por modelos de processamento de imagem (BOROVEC et al., 2017; OSCO et al., 2020; CHEN et al., 2020; AKBARI et al., 2020). Demonstrando potencial para análise visual dos filmes obtidos no processo industrial de recobrimento.

A exploração de sistemas de predição de qualidade do processo industrial de metalização por recobrimento utilizando modelos de ML e a interpretação das variáveis por técnicas de visualização fornecem informações importantes para o monitoramento e a manutenção do processo de manufatura de filme metálicos por recobrimento. Servindo de base para futuros trabalhos de implementação de sistemas de ML mais robustos e que possam ser integrados e fornecerem informações em tempo de produção.

1.1 OBJETIVOS

1.1.1 Objetivo geral

Este trabalho tem como objetivo utilizar dados operacionais e de qualidade de um processo industrial de produção de filmes metálicos por deposição para explorar aplicações de modelos de aprendizado de máquina (*Machine Learning*) (ML). Além de extrair correlações e interpretar comportamentos de variáveis relevantes para o processo de predição de qualidade dos produtos através de técnicas de visualização.

1.1.1.1 Objetivos específicos

- Levantamento de dados e informações do processo de produção e de verificação de qualidade do filme metalizado produzido.
- Cruzamento de informações para formação do base de dados.
- Construção de modelos para análise de qualidade do produto a partir de dados de processo.
- Construção de modelos de ML para predição de classe dos produtos.
- Aplicação de métodos de visualização para análise de produtos e de variáveis de processo.

1.2 CONTRIBUIÇÕES

Através deste trabalho foi possível o desenvolvimento de duas aplicações utilizando dados de processo e de qualidade envolvidos no processo de metalização a vácuo de filmes. A primeira que observa padrões e características de variáveis do processo que interferem na qualidade dos produtos obtidos, e indica produtos potencialmente reprováveis.

E a segunda aplicação, implementa técnicas de pré-processamento de imagem que permitem a interpretação dos produtos obtidos e a observação de possíveis falhas, consideradas na construção de um modelo de ML utilizado para predição de qualidade do filme.

A partir deste trabalho de dissertação foram desenvolvidos dois artigos, o primeiro com título *Visual Analysis of Deep Learning Methods for Industrial Vacuum Metalized Film Product*

que foi submetido para a conferência *Systems, Man and cybernetics Society 2021*. O segundo trabalho foi submetido para a revista *Expert Systems with Applications* da Elsevier, intitulado *Machine Learning and Visual Applications Analysis for Quality Prediction of Film Metalization Process*.

1.3 ORGANIZAÇÃO DO TRABALHO

O trabalho foi dividido em capítulos que apresentam a descrição do embasamento científico deste trabalho, com exemplos e explanação dos métodos aplicados para processamento de dados e de variáveis, além do processo de análise de qualidade e predição de classificação (capítulo 2). Já o capítulo 3 apresenta o procedimento metodológico para extração e processamento de dados e aplicação dos métodos de ML. No capítulo 4 são mostrados e discutidos os resultados obtidos. E por último, capítulo 5, foram observadas as contribuições da pesquisa desenvolvida e as possibilidades para trabalhos futuros.

2 FUNDAMENTAÇÃO TEÓRICA

2.1 PROCESSO INDUSTRIAL DE RECOBRIMENTO

2.1.1 Recobrimento de filmes poliméricos

Recobrimento é um termo geral que tem como principal significado a função de restrição total ou parcial a passagem de algum composto ou substância em um dado material ou ambiente, a fim de preservar a composição do sistema material-atmosfera (BISHOP, 2015). Cada aplicação de recobrimento busca uma restrição específica. Em aplicações alimentícias, o recobrimento por filmes metalizados objetiva a proteção de alimentos e derivados, restringindo a exposição do alimento a oxigênio, vapor de água e também da luz do sol, a fim de retardar a sua degradação. Outra aplicação bastante conhecida é a sobreposição metálica para criação de condutores plásticos, sendo a função do metal melhorar a condutividade do polímero ou liga a qual está sendo aplicado.

Os recobrimentos por deposição à vácuo possuem grande versatilidade, podendo ser aplicados a sistemas com rolos de filmes de até 4,45 m de largura com comprimentos em média de 45 km de extensão (Figura 1), ou em sistemas em escala reduzida, de centímetros de largura e menos que 100 m de comprimento. Nesse sistema de deposição, o metal é vaporizado utilizando barcas cerâmicas aquecidas por resistência e a deposição é feita em filmes com alta velocidade de enrolamento, que podem exceder 1250 m/min, ou para sistemas mais complexos, com menores velocidades de deposição e velocidades de enrolamento menores que 1 m/min, o que possibilita a produção de filmes recobertos de diversos padrões (BISHOP, 2011).

A performance do recobrimento é estabelecida de acordo com o material e o processo de aplicação, podendo ser utilizado em garrafas de vidro, latas e filmes poliméricos, com espessura de camada metálica variáveis e aplicação direta, como o processo por deposição à vácuo, ou a partir de sobreposição de uma lâmina metálica, para casos onde é necessário uma maior performance em aplicações críticas como o caso dos fármacos (BISHOP; III, 2016).

O processo de recobrimento por deposição de alumínio a vácuo (Figura 2) consiste basicamente em uma grande câmara de recobrimento a vácuo que contém uma superfície plana, na qual o filme polimérico, ou substrato, passa para ser recoberto pelo metal.

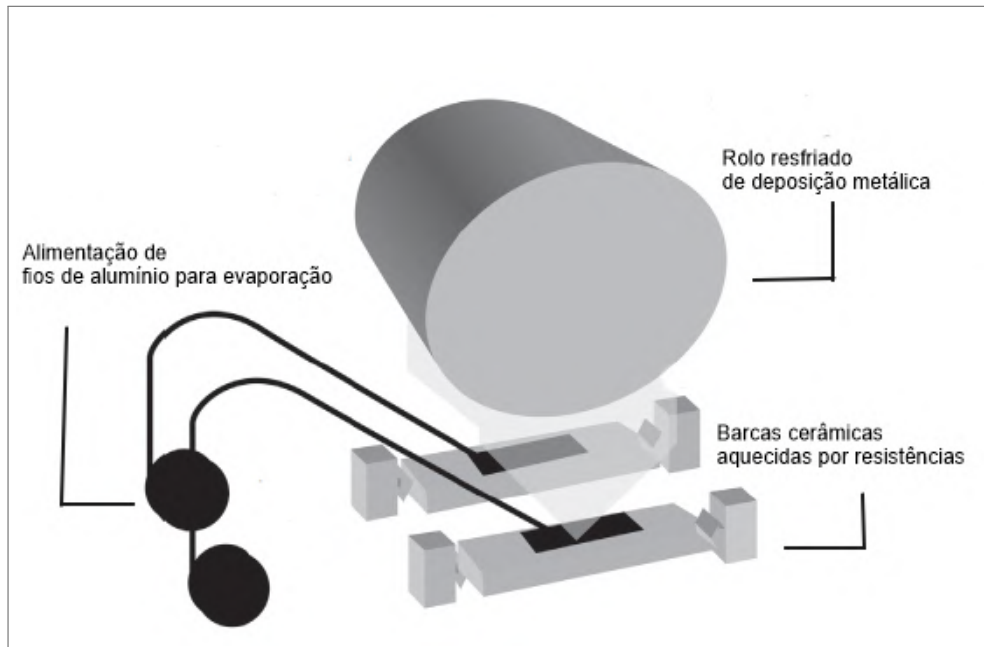
O material para deposição é alimentado na forma de fios e empurrados através de rolos que o comprimem e empurram contra uma barca cerâmica aquecida através de uma resistência

Figura 1 – Imagem do produto final do processo industrial de metalização a vácuo. O filme metalizado pode apresentar dimensões variadas, com até 4,45 m de largura e até 85 km de extensão.



Fonte: <<https://bit.ly/2MwNH0f>>

Figura 2 – Ilustração da zona de deposição presente no processo de metalização a vácuo. Destacando as barcas cerâmicas dispostas transversalmente em relação ao comprimento do filme, e a alimentação de fios metálicos por bobinas de passo para liquefação e vaporização do metal pelas barcas aquecidas. Por fim, o metal vaporizado é depositado sobre a superfície do filme.



Fonte: BISHOP (2011) adaptado.

elétrica, de forma que ao fio metálico entrar em contato com a barca, forme um pequeno lago de material fundido e evapore para formação de uma nuvem de vapor metálico para deposição no substrato, sobre um rolo de deposição resfriado, que passa pela zona de deposição. Como forma de reduzir reações químicas de oxidação ou adesão de partículas externas ao aerossol

metálico formado, o processo ocorre a vácuo (BISHOP, 2007).

2.1.2 Desafios do processo de recobrimento metálico de filmes por deposição

Apesar do processo industrial de recobrimento já estar presente há décadas dentre os processos de manufatura, ainda são encontradas dificuldades e gargalos no seu processo produtivo. O primeiro dos desafios é a garantia de uniformidade do recobrimento, uma vez que o processo de deposição ocorre em alta velocidade e que apresenta alta sensibilidade a diversas condições de processamento, como: temperatura das barcas cerâmicas, velocidade do fio de alumínio, a limpeza da superfície do substrato para receber o alumínio vaporizado, e o controle das demais variáveis operacionais.

Na atualidade, o processo industrial de recobrimento analisado nesse trabalho, apresenta controle automático apenas das taxas de evaporação do metal, ou seja, é observada apenas a quantidade de metal a ser depositado através das medidas de densidade óptica (*Optical Density*) (OD). Essa análise limitada do processo e a falta de conhecimento das variáveis e da condições operacionais ideais, levam muitas vezes a operações instáveis e a formação de recobrimentos desuniformes, reduzindo a qualidade dos produtos e a otimização do uso de recursos.

Por outro lado, também existem limitações na análise dos produtos formados, uma vez que se tratam de produtos com grandes dimensões, o que limita a análise dos filmes produzidos à inspeção de pequenos recortes desses produtos para análise por especialistas e por equipamentos de laboratório. Essa limitação de área inspecionada pode fornecer a obtenção de resultados não representativos da qualidade geral do filme. E mesmo as análises laboratoriais feitas em pequenos recortes do filme representam também gargalos para a produção, uma vez que análises como a de permeabilidade a vapor e permeabilidade a oxigênio demandam de oito a 12 horas, retardando o processo de liberação desses produtos para entrega.

A falta de inspeções mais robustas e da análise de parâmetros operacionais de forma mais completa contribuem para a obtenção de taxas elevadas de produtos reprovados, e ainda, uma análise incompleta da qualidade dos produtos formados. Assim, esse trabalho busca solucionar ou contribuir para a redução dessas dificuldades, dando suporte de informação e fazendo previsões de variáveis através de modelos de ML.

2.1.3 Variáveis de processo e de qualidade

Sabe-se que a qualidade do processo de recobrimento metálico apresenta diversas variáveis bastante sensíveis que podem ou não se correlacionar, e que a primeira delas é a qualidade do filme a ser recoberto. Isto inclui aditivos ou produtos residuais presentes na superfície do filme resultantes do processo de polimerização utilizado para a sua formação, que pode resultar em uma redução da energia superficial do substrato e causar uma baixa energia de adesão do metal, causando grandes interferências na forma de preenchimento da superfície do filme pelo metal (Figura 3), demonstrando que filmes formados por estruturas poliméricas diferentes também podem causar variações no processo. Como forma de solução desse problema existem diversas formas de tratamento de superfície do substrato ou do ambiente, uma delas é o uso de plasma (BISHOP; Mount III, 2012).

Figura 3 – Efeito da energia superficial sobre a qualidade do recobrimento. A depender das condições da superfície do filme para metalização, pode ocorrer a formação de filmes com altos valores de densidade óptica, porém com recobrimento pobre.



Fonte: BISHOP; III (2016) adaptado.

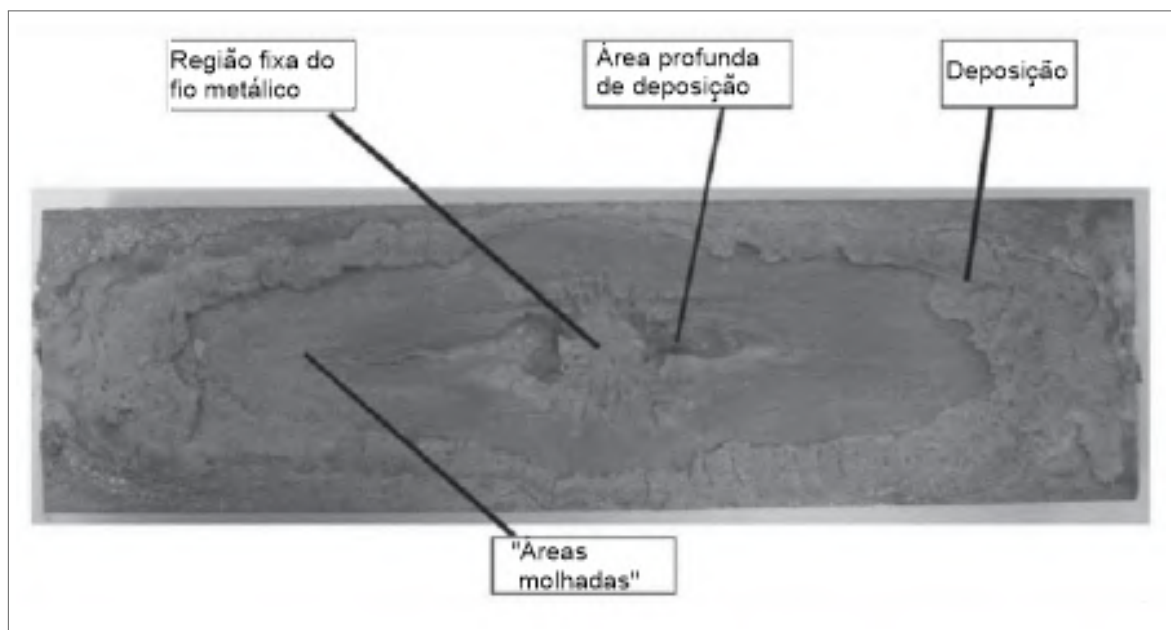
Outra variável bastante importante na qualidade do produto a ser processado é o substrato metálico depositado, os fios metálicos podem apresentar variações de diâmetro, dureza, pureza e limpeza (Mount III, 2008), variáveis essas que influenciam diretamente na redução de furos no recobrimento e consequente melhoramento na performance da barreira formada.

A forma como o fio do substrato metálico é alimentado também causa alterações no processo, pois, com a sua evaporação o tubo condutor do fio metálico pode ser parcialmente ou totalmente obstruído, causando o desvio do fio metálico em direção a barca evaporadora aquecida, o que leva a problemas na formação do lago de metal fundido na barca evaporadora, necessário para evaporação correta, ou até a interrupção da alimentação do fio metálico, causando a má/não formação da nuvem metálica. O formato do lago formado sobre a barca é também correlacionado diretamente a temperatura em que a barca se encontra, pois, a taxa

de evaporação do metal é proporcional ao calor fornecido a barca, o que remete diretamente a diferença de potencial aplicada a resistência elétrica e a condição presente da barca para seu aquecimento.

As barcas são compostas principalmente por uma mistura de duas fases intermetálicas de Nitreto de Boro (BN) e Diboreto de Titânio (TiB_2) que inicialmente são igualmente distribuídas e ligadas, porém, com a sua utilização, ocorre um processo de migração e formação de diferentes regiões com concentrações desbalanceadas entre as fases, o que causa variações na resistividade do material e na uniformidade de superfície para onde o lago de alumínio é formado, levando a variações na qualidade da cobertura metálica do substrato (Figura 4). Algumas barcas também possuem uma terceira fase composta por Nitreto de Alumínio (AlN), com o objetivo de obter a melhor combinação de resistividade e resistência a corrosão (BISHOP, 2015).

Figura 4 – Ocorrência de áreas de diferentes aspectos causadas pelo desgaste em uma barca de evaporação metálica. Essa desuniformidade na superfície cerâmica pode provocar o surgimento de falhas no recobrimento do filme na sua respectiva posição, levando a formação de produtos de baixa qualidade.



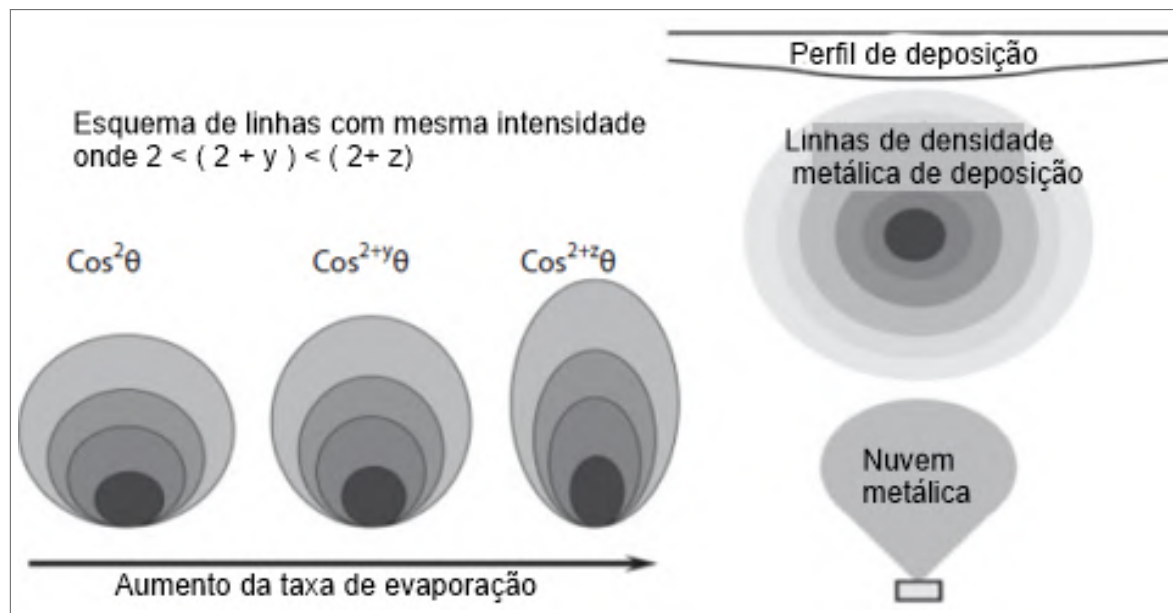
Fonte: BISHOP (2015) adaptado.

Normalmente, as barcas são instaladas todas com os mesmos valores de resistividade para que possuam o mesmo tempo útil e os mesmos valores de resistência. Barcas de qualidade elevadas possuem tempo de vida útil esperado de 15 ± 5 h, ou 90 ± 30 min, porém esses valores variam bastante de acordo como as barcas são utilizadas (Mount III, 2008).

Conforme o tempo de uso das barcas, a voltagem tende a reduzir 60%-70%, mudança

essa causada pela erosão e por alterações químicas sofridas pelas barcas devido a corrosão e a reações químicas, fazendo com que a corrente elétrica aplicada as barcas seja aumentada em torno de 40% - 50% para que a temperatura do evaporadores seja mantida. A Figura 5 demonstra o efeito da variação de temperatura e a sua importância para formação da nuvem de vapor metálico que afeta diretamente na uniformidade do recobrimento.

Figura 5 – Dependência do perfil da nuvem metálica de deposição de acordo com a temperatura da barca. Esse feito demonstra a importância direta da estabilidade da temperatura presente nas barcas cerâmicas, além da necessidade que todo o conjunto de barcas apresente performance semelhante de estabilidade de temperatura, de forma a manter uniformidade no recobrimento metálico do filme.



Fonte: BISHOP (2015) adaptado.

De acordo com Archibald e Parent (1976) essa mudança de performance pode ser observada em função da taxa de alimentação do fio metálico ou da temperatura presente nos evaporadores, e relacionadas as dimensões da barca, a sua resistividade, voltagem inicial, corrente ou energia.

Todas essas variáveis se combinam com a intenção de obter um recobrimento metálico mais uniforme e com valores de OD próximos ao pré-determinado para o produto. Uma vez que essa variável mede a capacidade de barreira do material contra a passagem de luz, dada

pela Equação 2.1 (BISHOP; III, 2016).

$$\begin{aligned}
 \text{Opacidade} &= \frac{\text{Luz incidida}}{\text{Luz transmitida}} \\
 \text{Transmitancia (T)} &= \frac{\text{Luz transmitida}}{\text{Luz incidida}} \\
 \text{Opacidade} &= \frac{1}{T} \\
 \text{Densidade optica (OD)} &= \log_{10} \text{Opacidade} = \log_{10} \frac{1}{T}
 \end{aligned} \tag{2.1}$$

Assim, durante o processo de metalização as taxas de evaporação do metal são controladas de acordo com os valores de OD obtidos em tempo de processo, através das variáveis de velocidade do fio e de temperatura da barca. Fazendo com que a variável OD represente o ponto de controle principal do processo industrial de recobrimento metálico de filmes.

2.1.4 Processo industrial de deposição metálica em filme

Através de uma parceria com uma indústria, foi analisado o processo industrial de recobrimento metálico de filme através da metalizadora TOPMET 2450 (Figura 6). Esse processo industrial ocorre em bateladas, quando um produto é formado por vez e de forma pausada. O tempo para a recobrimento de um filme varia de acordo com as dimensões do filme a ser recoberto, também conhecido como substrato, que podem apresentar largura de 2,5 a 4,5 m e comprimento de até 87 km. Ao fim de uma batelada, um novo filme é carregado na máquina na forma de rolo, o qual é desenrolado, metalizado e enrolado na forma de um novo rolo (BISHOP, 2007; PERRY; LENTZ, 2009).

A metalizadora TOPMET 2450 conta com vinte e seis evaporadores e, alinhados a esses, vinte e seis sensores, posicionados no sentido transversal ao filmes e distantes 0,1 m entre si, para a medição da transmitância do filme metalizado. Ao início de uma batelada é determinado o valor desejado de OD, ou transmitância, e variáveis de pressão interna, tempo de ventilação e temperatura do tambor de resfriamento. A resistência aplicada aos evaporadores e a velocidade de alimentação dos fios de alumínio são ajustadas automaticamente e de forma individual, influenciando no recobrimento da posição da barca evaporadora e nas duas posições vizinhas mais próximas (TOPMET, 2016).

A alimentação dos fios metálicos que são fundidos e evaporados sobre o filme é feita por bombas com motores de passos, com cerca de 400 posições por giro, com precisão de até

Figura 6 – Imagem da metalizadora TopMet2450 utilizada para processo industrial de recobrimento metálico em filme polimérico.



Fonte: TOPMET (2016) adaptado.

$\pm 0,02\%$ no controle da velocidade, o que garante estabilidade na deposição metálica. O equipamento produz um alto vácuo na câmara de deposição na ordem de 10^{-6} a 10^{-5} mbar, em seguida aquece as barcas evaporadores através de resistências até temperatura superior ao ponto de fusão do metal a ser depositado, temperatura essa na ordem de 1400°C para a evaporação do alumínio (BISHOP, 2007).

Ao início da evaporação do metal, o vácuo da câmara de deposição passa a ser da ordem de 10^{-4} mbar e o filme é desenrolado com velocidade média de 400 m/min para ser recoberto. Durante o processo de deposição, os valores de OD são mostrados em tempo real em uma tela acoplada ao equipamento, em uma faixa de azul a vermelho, onde o azul intenso demonstra valores de OD alta, o que representa uma quantidade elevada de recobrimento, e a cor vermelha representa uma OD baixa, o que representa um baixo recobrimento, ambos em relação ao valor desejado e determinado antes de iniciar a batelada (BISHOP, 2007; PERRY; LENTZ, 2009).

Em tempo de processo, são armazenados os resultados de OD a cada dois minutos associado ao comprimento do filme em que foi feita a medição, além disso, são registradas as variáveis com valores de referência e com valores operacionais presentes no Quadro 1, junto com as variáveis de qualidade.

Outras variáveis de produção são armazenadas em um programa de armazenamento de dados, TOTVs®, variáveis essas como: bobina, campanha, máquina, data de produção, quantidade de furos, arranhão, e estatísticas de OD, como, Desvio Padrão OD, Média OD e OD 1 a OD 11, que servem para registro de informações e cruzamento de dados.

Como visto, o processo de recobrimento metálico envolve diversas variáveis que muitas vezes estão correlacionadas entre si, parâmetros como tipo de filme, pureza do fio metálico,

Quadro 1 – Variáveis envolvidas no processo de metalização à vácuo de filme

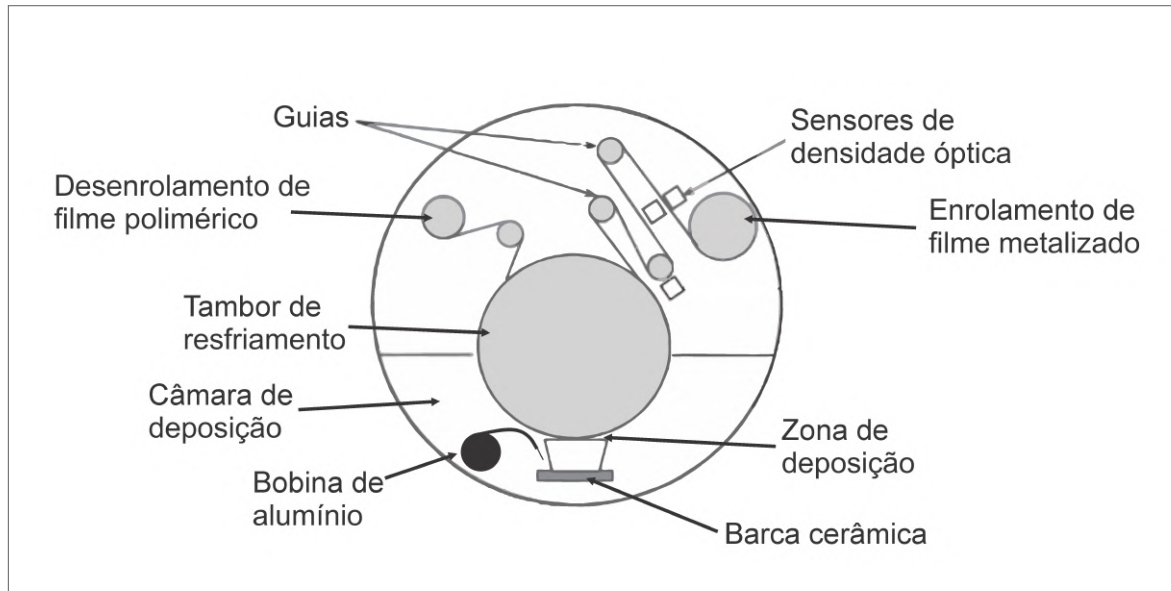
Tipos de variáveis	Variáveis
Valores de referência	Densidade óptica (OD), valor de referência, tensão inferior, tensão de linha, limite superior e limite inferior
Valores operacionais	Tempo de aquecimento, tempo de recobrimento, temperatura do tambor de resfriamento, tempo de vácuo, tempo de ventilação, tempo de vida da barca, velocidade do filme, velocidade do fio de alumínio
Valores de qualidade laboratorial	Espessura, permeabilidade a água, permeabilidade a vapor, densidade óptica média, remoção de alumínio, classificação do filme.

Fonte: Elaborada pelo autor (2021)

tempo de uso das barcas, resistência e pressão na câmara de deposição, cada uma dessas relacionadas a uma etapa do processo de produção do filme, conforme ilustrado na Figura 7, determinante para a qualidade do produto. Além dos valores furos (*Pinholes*), falhas maiores presentes na superfície do filme são caracterizadas como arranhão, e a espessura da barreira formada pelo recobrimento é medida a partir de sensores de OD distribuídos de forma transversal ao filme e alinhados as barcas intermetálicas, esse conjunto de variáveis reproduzem a qualidade do produto a partir de dados coletados no processo (PERRY; LENTZ, 2009).

De forma complementar, recortes de filmes recobertos são retirados dos rolos e levadas ao laboratório de análise de qualidade, onde são avaliadas as variáveis de qualidade mostradas no Quadro 1. Esse conjunto de dados é base para caracterização da qualidade do filme em laboratório, porém, até o presente momento, apenas essas variáveis analisadas são levadas em consideração para emissão ou não de alerta de qualidade inferior a desejada, e para a classificação final do filme como aprovado ou reprovado.

Figura 7 – Ilustração do processo de metalização à vácuo. O processo inicia com a alimentação de um filme polimérico (substrato) a ser desenrolado, guiado até a zona de deposição e recoberto pelo metal vaporizado em barcas cerâmicas. Esse metal é alimentado em forma de fios por bobinas de passo, para serem liquefeitos, vaporizados e depositados sobre a superfície do filme. Após o recobrimento, através de sensores são obtidas informações de qualidade do recobrimento. E por fim o filme recoberto é enrolado novamente.



Fonte: BISHOP (2011) adaptado.

2.2 INTELIGÊNCIA ARTIFICIAL

A inteligência artificial (*Artificial Intelligence*) (AI) tem como uma de suas principais funções a construção de sistemas que permitem através de algoritmos a realização de uma ou mais atividades como: compreensão de linguagem natural, desenvolver atividades mecânicas complexas, solucionar problemas computacionais complexos que possam envolver grandes quantidades de dados e oferecer respostas compreensíveis ao ser humano (JOSHI, 2020). A relevância de implementações envolvendo AI para o setor industrial é dada pela possibilidade do processamento e obtenção de informações rápidas e de forma abrangente para aplicação em processos. Sendo por isso, atualmente, tida como uma das principais ferramentas para obtenção de informações, conhecimentos de negócios e apoio a decisões operacionais (DAOUTIDIS et al., 2018).

Além disso, os algoritmos de AI possuem aplicações para otimização e controle operacional dada a sua capacidade adaptativa a mudanças rápidas, podendo entregar soluções simultâneas para produtividade, sustentabilidade, flexibilidade e falhas no processo, através de várias técnicas de otimização matemática e de modelos de ML (SHIN; LEE; REALFF, 2017; SHIN et al., 2019).

As aplicações de ML presentes na literatura se referem majoritariamente ao uso de modelos de ML específicos e ao seu ajuste de hiper-parâmetros para obtenção de soluções com melhores performances de predição. Diferentemente desse trabalho, que se propõe a comparação de forma compreensível de diferentes técnicas de ML para modelagem das variáveis do processo industrial de recobrimento metálico de filmes. Uma vez que os processos de implementação desses modelos são similares, e a sua comparação serve de base para seleção de modelos que melhores se adaptem ao problema e que forneçam sistemas mais robustos (PATEL et al., 2010).

Comparar os diferentes sistemas de ML consiste no desenvolvimento de um espaço de exploração, que envolve o uso de diferentes técnicas de pre-processamento de dados, de pre-processamento de variáveis e de ajustes de parâmetros e hiper-parâmetros para diferentes algoritmos de ML (WEI; ZHAO; MIAO, 2018). E para além disso, avaliar os sistemas desenvolvidos pelas suas performances e compará-los com a importância das variáveis para o desempenho real do processo industrial (OZISIKYILMAZ; MEMIK; CHOUDHARY, 2008).

Assim, esse trabalho se preocupa na exploração dos modelos de ML disponíveis na literatura para formarem predições sobre a qualidade final dos produtos obtidos no processo de recobrimento metálico por deposição. Para além disso, também são analisadas as performances desses modelos e como as suas predições são construídas. Servindo de base para a escolha de um modelo de ML adequado para ser posto em produção através de trabalhos futuros.

2.2.1 Algoritmos de aprendizagem de máquina

aprendizado de máquina (*Machine Learning*) (ML), ou do inglês *Machine Learning*, é um ramo de estudo da AI no qual um modelo ou programa computacional consegue aprender, gerar correlações e executar atividades não explicitamente programadas baseando-se em dados representativos de um sistema (DANGETI, 2017). Para aplicações industriais, os modelos são inicialmente utilizados para análise das variáveis, e em seguida, para exploração dos sistemas de ML possíveis para melhoria de qualidade do processo sem interferir nos parâmetros operacionais. Após isso, essas aplicações buscam o desenvolvimento de sistemas em produção, para interação com os parâmetros de manufatura. E por fim, é estabelecido um sistema integrado para análise de dados, de modelos de ML e otimização de parâmetros operacionais (WEICHERT et al., 2019).

Para esse trabalho, será desenvolvida a etapa de exploração de sistemas sem diretamente interferir no processo produtivo. E a exploração e execução de algoritmos de ML com essa

finalidade é baseada, simplificadamente, em três fatores, os dados utilizados, as métricas escolhidas, e as respostas obtidas.

Esses modelos são classificados de acordo com o seu tipo de aprendizagem podendo ser considerados supervisionados, não-supervisionados e por reforço (Figura 8). A diferenciação das classes é feita de acordo com o tipo dos dados utilizados, sendo para conjuntos de dados rotulados utilizados os algoritmos de aprendizado supervisionado, e para dados não-rotulados os algoritmos são conhecidos como não-supervisionados (JOSHI, 2020). Outra categoria a parte dessas, é o aprendizado por reforço, que considera para tomada de decisões as observações de impacto sobre o ambiente a cada ação tomada (HAO; ZHANG; MA, 2016).

Figura 8 – Subdivisão dos modelos de aprendizagem de máquina de acordo com o seu tipo de aprendizado e suas principais aplicações.



Fonte: Elaborada pelo autor (2021)

A forma como são desenvolvidos os algoritmos de ML, baseados em conceitos matemáticos e estatísticos, oferecem uma gama de aplicações em áreas da engenharia, processamento de sinais, análises financeiras, ciência genética e entre outras (JOSHI, 2020). Cada forma de aprendizagem dos modelos de ML apresenta um conjunto de aplicações específicas, de forma a otimizar os principais objetivos dessas aplicações, que são a generalização e discriminação desses eventos.

Algoritmos de aprendizado supervisionado tem em sua estrutura de dados, um conjunto

de variáveis associados a um conjunto de saídas ou rótulos, no qual o objetivo do algoritmo de ML, basicamente, é obter as correlações entre as variáveis que conduzem a predição de valores ou de classes dos objetos que constituem o banco de dados. Apesar da dificuldade de se obter bancos de dados rotulados, existem diversas aplicações que utilizam desses algoritmos para obtenção de informação, como os problemas de reconhecimento de imagem, diagnóstico de sistemas e produtos e também previsões de mercado.

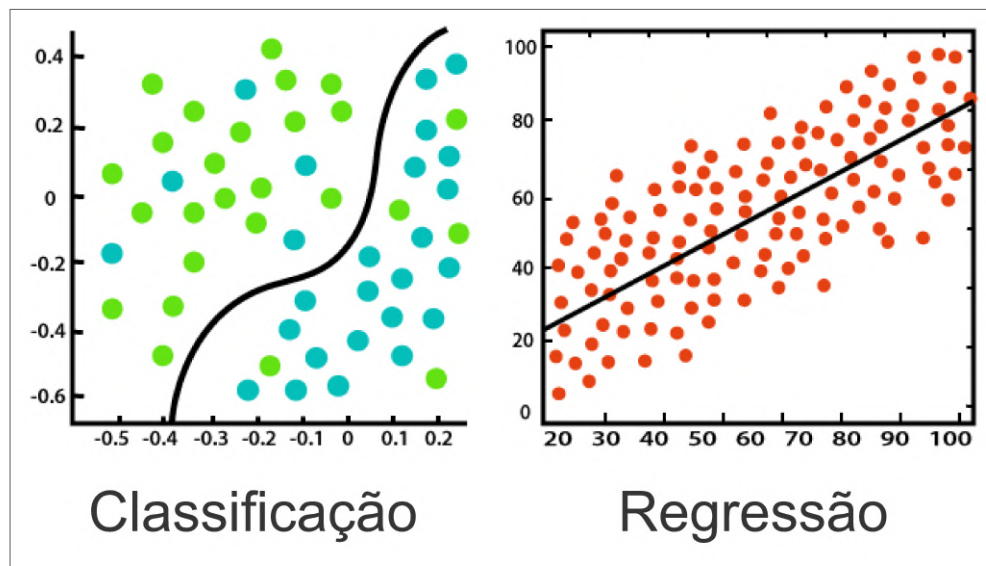
Diferentemente dos algoritmos de aprendizado não-supervisionado, nos quais não existem classes pré-definidas e o objetivo do algoritmo é a extração de padrões e classes dentro de um conjunto de informações, em geral, essa técnica é utilizada como forma exploratória dos dados e aplicadas para elicitación de atributos, sistemas de recomendação e segmentação de clientes. Dentre os principais métodos não-supervisionados existem: K-means (MACQUEEN et al., 1967), análise por componentes principais (*Principal Component Analysis*) (PCA) (BISHOP, 1999), aproximação e projeção uniforme de múltiplas dobras (*Uniform Manifold Approximation and Projection*) (UMAP) (MCINNES; HEALY; MELVILLE, 2020), incorporação estocástico de vizinho distribuído em T (*T-distributed Stochastic Neighbor Embedding*) (TSNE) (MAATEN; HINTON, 2008), *Singular Value Decoposition* (HENRY; HOFRICHTER, 1992) e *Deep Auto Encoders* (KRAMER, 1991).

Por outro lado, o aprendizado por reforço não exige necessariamente a construção prévia de um banco de dados, pois o seu método explora ações em um sistema e as seleciona de forma a maximizar um função pré-determinada que associa a condição atual do problema com a condição desejada, conhecida como métrica de recompensa, passando assim para um momento posterior de exploração de ação no sistema em um novo estado, a análise de recompensa, e por fim, a tomada de novas ações de forma subsequente (SUTTON; BARTO, 2018a). De forma simplificada, o seu desempenho é otimizado de acordo com ações, respostas e o seu processamento, obtidas para uma dada atividade ao longo do tempo, com base na interação contínua com o ambiente, geralmente desconhecido (SUTTON; BARTO, 2018b). Essa interação pode ser feita a partir de 3 formas de sinais: (1) sinal de estado, que descreve o estado do processo; (2) sinal de ação, que permite influenciar o processo; (3) sinal de recompensa, que retorna para o agente uma resposta sobre seu desempenho imediato. Esse tipo de algoritmo é utilizado, principalmente, em ambientes que exigem tomadas constantes de ação e aprendizagem, como as atividades de controle de processos industriais em tempo real.

2.2.1.1 Aprendizagem supervisionada

Os modelos de ML supervisionados, além de um conjunto de variáveis que caracterizam o problema, apresentam um conjunto de rótulos, que podem ser classes ou valores associados a cada objeto do banco de dados. A aprendizagem supervisionada apresenta dois segmentos: os modelos de classificação e os modelos de regressão. Os problemas de classificação tem como objetivo determinar a qual classe um determinado objeto não-rotulado pertence, utilizando para isso objetos rotulados. Já os problemas supervisionados voltados para regressão, possuem como rótulo valores numéricos, e o objetivo do modelo de ML construído é de acordo com o comportamento dos objetos da base de treino, modelar o problema para que possa atribuir rótulos a novos objetos (AMR, 2020) (Figura 9).

Figura 9 – Tipos de predição para modelos supervisionados de aprendizagem de máquina. A predição de classe tem como objetivo a classificação de novos objetos em um espaço com pelo menos duas classes representadas. Já a regressão é feita através da modelagem de exemplos pertencentes a mesma classe para obtenção de valores de variáveis para novos objetos.



Fonte: BELKACEM (2021) adaptado.

Para dados não supervisionados que se queira rotular, modelos de ML que buscam gerar agrupamento, e a partir disso, rótulos para os dados, como o k-vizinhos mais próximos (*K-Nearest Neighbors*) (KNN), ou k-vizinhos mais próximos, modelo esse que não apresenta equações ou relação entre entradas e saídas, mas que se baseia na distribuição espacial dos exemplos de treinamento e em métricas de distância (JOSHI, 2020). Ou as redes neurais artificiais (*Artificial Neural Networks*) (ANN), que são construídas a partir de múltiplos perceptrons, esses descritos pela multiplicação do vetor de variáveis n-dimensional x e o vetor de coeficientes

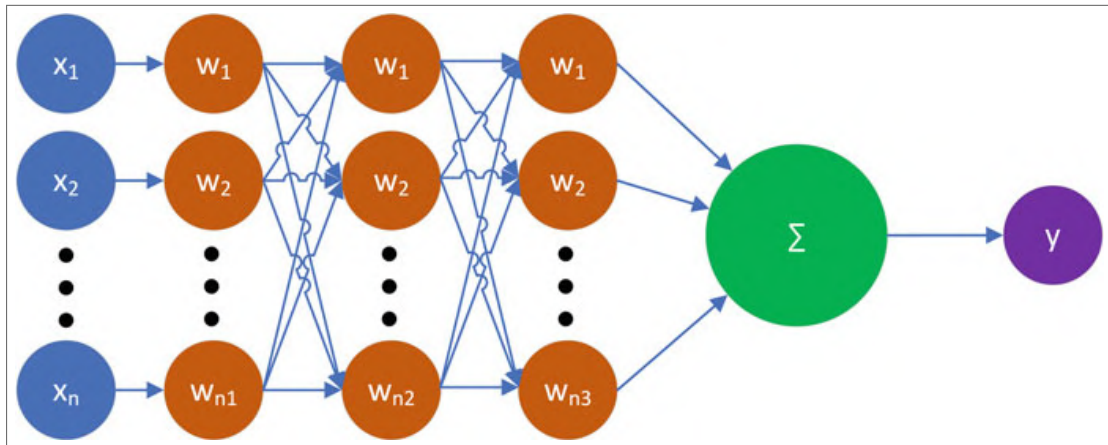
ou pesos w escrito na forma:

$$x.w = y \quad (2.2)$$

Os perceptrons nas ANN são interconectados e podem ser dispostos em múltiplas camadas sequenciais com unidades paralelas iguais a da Figura 10. Essa combinação de perceptrons em paralelo e em sequencia formam as conhecidas ANN, sendo o perceptron multicamadas (*Multi-Layer Perceptron*) (MLP) o mais básico dentre esses, já que seus perceptron são associados a funções de ativação lineares simples, também conhecidas como perceptrons.

Essa função de ativação tem como entrada os valores y dos perceptrons e como saída os valores zero ou um a depender de um valor de limiar. As combinações desses perceptrons podem ser utilizadas para diversos propósitos, como agrupamento, redução de dimensionalidade, regressão e classificação de objetos.

Figura 10 – Ilustração de um perceptron, estrutura básica do algoritmo de aprendizado profundo baseado em redes neurais, o perceptron multicamadas (*Multi-Layer Perceptron*) (MLP). As camadas de entrada são as variáveis X_i e o Y são os rótulos preditos pelo modelo. Já os parâmetros W_i e a função de ativação Σ são, respectivamente, os parâmetros e o hiper-parâmetros do modelo.

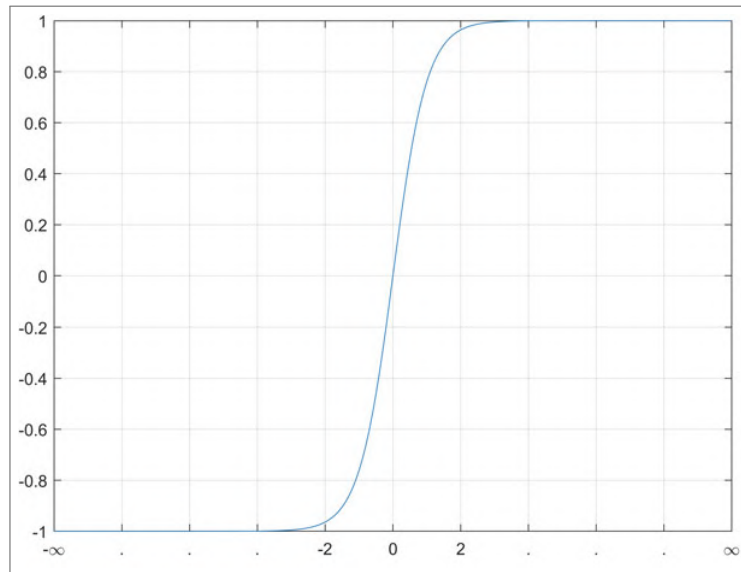


Fonte: JOSHI (2020)

Para casos não-lineares, as ANN utilizam de funções de ativação não-lineares que utilizam as saídas y como entrada para manipulação das informações de acordo com uma função matemáticas não-lineares de transformação (JOSHI, 2020). A Figura 11, mostra o exemplo onde a função de ativação é uma tangente hiperbólica, que recebe o valor de y e entrega como resultado o valor de $\tanh y$.

Outro exemplo bastante conhecido de função de ativação não-linear é a unidade linear retificada (*Rectified Linear Unit*) (ReLU), cuja função é dada pela equação $f(x) = \max(0, x)$,

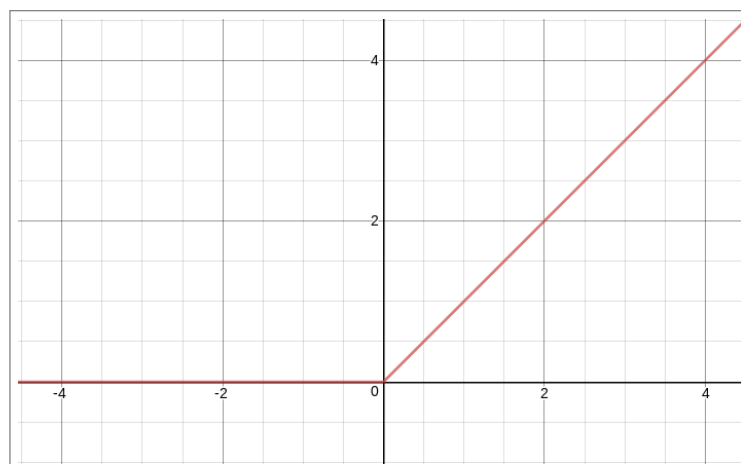
Figura 11 – Função de ativação tangente hiperbólica, $\tanh$, utilizada em modelos de ANN para modelagem de problemas não-lineares. A não-linearidade dos modelos é obtida através da transformação dos valores y fornecidos pelos perceptrons em valores de $\tanh y$.



Fonte: JOSHI (2020).

e mostrada na Figura 12 mostra o comportamento dessa função que quando recebe o valor de y negativo retorna zero, e quando y é positivo retorna uma função linear com inclinação igual a um. Existem diversos outros modelos de função de ativação (NWANKPA et al., 2018), o que torna a sua escolha dependente do problema que se deseja modelar, sendo por isso um dos hiper-parâmetros presentes no desenvolvimento de aplicação dos modelos de ANN.

Figura 12 – Função de ativação unidade linear retificada (*Rectified Linear Unit*) (ReLU) utilizada em modelos de ANN para modelagem de problemas não-lineares. Para valores negativos a função retorna zero, e valores positivos a função tem comportamento linear com inclinação igual a um.



Fonte: AGARAP (2018).

Outro modelo supervisionado bastante utilizado na literatura é o árvore de decisão, que

é construído por regras encadeadas geradas para aproximação de funções discretas capaz de fazer inferências mesmo através de funções disjuntas (MITCHELL, 1997). Outras vantagens importantes nos modelos de árvore de decisão são a possibilidade de manipulação de dados não-numéricos, a sua capacidade interpretativa e sua escalabilidade.

Desenvolvido por Corinna Cortes e Vladimir Vapnik 1995, o modelo de máquina de Vetores de Suporte (*Support Vector Machine*) (SVM) se baseia na construção de um hiperplano para separação de classes, sendo o hiperplano construído através de exemplos de fronteira entre as classes, o que reduz o gasto computacional para armazenamento de dados de treinamento (JOSHI, 2020). Além dos tipos de modelos supervisionados mostrados, existem modelos mais recentes como os algoritmos evolucionários e os algoritmos baseados em aprendizado profundo, que apresentam maior complexidade, mas também maior capacidade de adaptação a problemas não lineares (JOSHI, 2020).

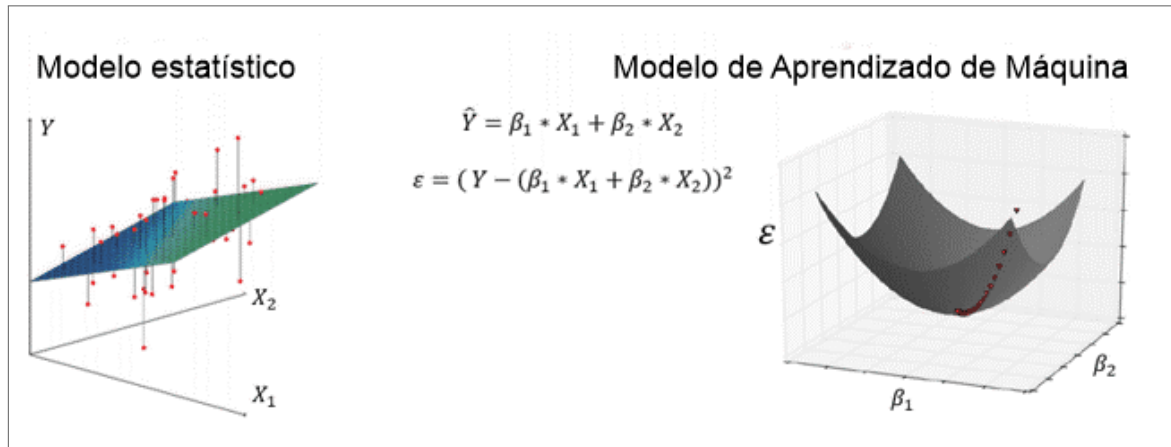
Modelos supervisionados de ML são amplamente utilizados para regressão de variáveis em diversas áreas (NING; YOU, 2019), sendo a indústria um dos principais setores, com aplicações em monitoramento de processo, diagnóstico de falhas e desenvolvimento de sensores inteligentes (*soft-sensors*) (Ge; Chen, 2019). Por décadas, essas áreas eram desenvolvidas, basicamente, através de modelos estatísticos como os modelo de Controle Preditivo (*Model Predictive Control*) (MPC), porém, a capacidade adaptativa dos modelos de regressão de ML tem despertado interesse da comunidade acadêmica e industrial devido a sua natureza adaptativa (DAOUTIDIS et al., 2018).

2.2.1.2 Modelos estatísticos x modelos de ML

Modelos estatísticos de regressão e modelos computacionais de ML tem como objetivo a modelagem de um sistema de acordo com um conjunto de exemplos que compõem o seu espaço amostral, porém através de formas de diferentes de execução. Modelos estatísticos buscam a criação de um hiperplano com função matemática conhecida de forma a minimizar a soma dos erros entre os valores fornecidos pelo hiperplano e as observações verdadeiras. Já para os modelos de ML o objetivo é, através dos dados, construir o modelo estabelecendo os seus parâmetros e hiper-parâmetros de forma a minimizar a soma dos erros, ϵ , de predição de uma variável externa, seja ela valores ou classes (Figura 13).

A natureza paramétrica dos modelos estatísticos de regressão, o que significa que o modelo é construído por parâmetros que são diagnosticados para validade do modelo, os diferenciam

Figura 13 – Diferenciação do processo de modelagem por modelos estatísticos e por modelos de aprendizado de máquina (*Machine Learning*) (ML). Os modelos estatísticos utilizam de funções matemáticas para representação das variáveis, cujo o objetivo é a redução do erro entre a representação dos exemplos e os próprios exemplos. Já os modelos de ML tem como função objetivo a minimização do erro entre valores de predição e os valores verdadeiros atribuídos aos objetos de um espaço, utilizando para isso o ajuste de parâmetros e hiper-parâmetros construídos e fornecidos, respectivamente.



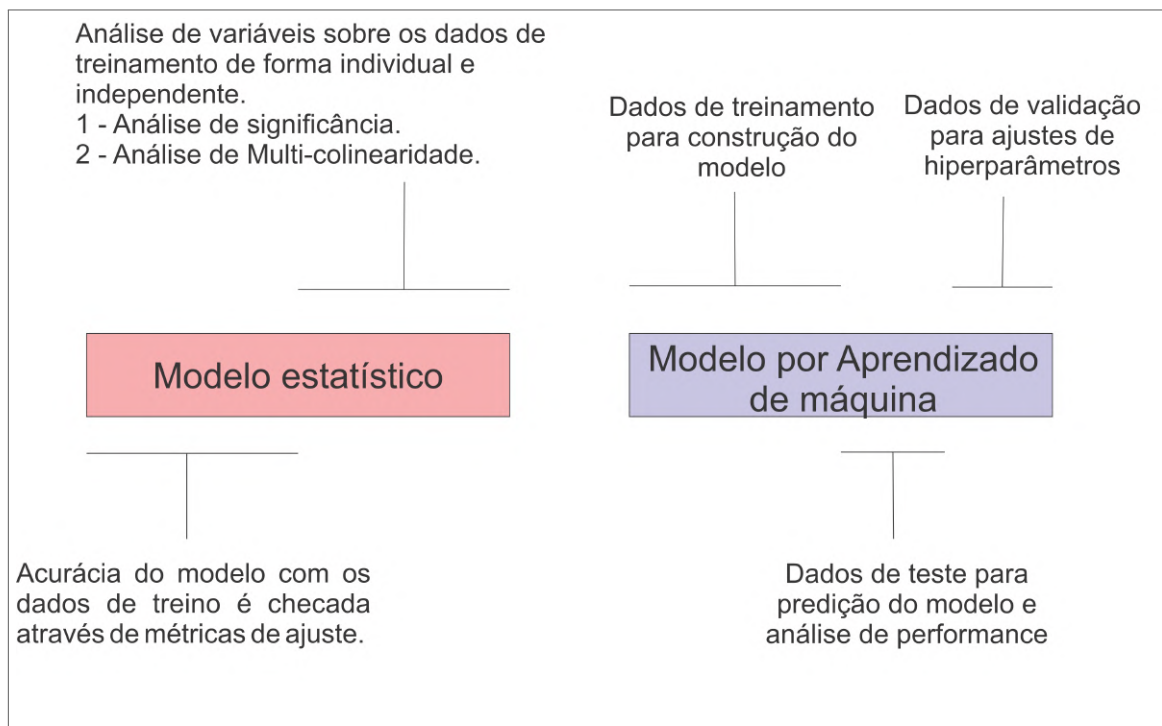
Fonte: BISHOP (2015) adaptado.

dos modelos de ML, que não apresentam suposição prévia de curva de parametrização, o que permite a esses modelos, o treinamento de forma direta através dos dados fornecidos, propondo modelos complexos de função com melhores ajustes ao invés de modelos pré-definidos (DANGETI, 2017).

Outra diferença apresentada pelos modelos de ML é o particionamento dos dados que são divididos em dados de treino, de validação e de teste, devido a grande liberdade na formação dos seus modelos. Diferentemente dos modelos estatísticos que devido a utilização de funções matemáticas pré-estabelecidas, em geral, os dados são particionados para análise de multicolinearidade e para a definição dos parâmetros do modelo, para os algoritmos de ML os dados de treinamento servem para a construção do modelo, que é avaliado com os dados de teste e os dados de validação servem para ajuste de hiper-parâmetros (Figura 14) (DANGETI, 2017). Isso fornece a vantagem aos modelos de ML, pois, dessa forma apresentam duas formas de validação do modelo para garantia de robustez em suas predições.

O ajuste automático dos pesos, os métodos de otimização com base em dados e a não necessidade de parametrização de problemas apresentados pelos modelos de ML fazem com que o seu uso represente importantes avanços na construção de modelos mais precisos e de maior performance que os modelos estatísticos tradicionais, sendo por isso, tidos como recursos promissores em aplicações de processamento de dados em indústrias de processamento (BIKMUHAMETOV; JäSCHKE, 2020).

Figura 14 – Modelos estatísticos utilizam, em geral, uma subdivisão dos dados para análise de multicolinearidade e outra parte para análise de ajuste de funções paramétricas. Por outro lado, os modelos de aprendizado de máquina são construídos com a subdivisão de treinamento, ajustados com a subdivisão de validação e verificados com a subdivisão de teste, o que fornece mais robustez para generalização do modelos com novos exemplos.



Fonte: Elaborada pelo autor (2021)

2.2.1.3 Algoritmos de ML e os processos industriais

Radhakrishnan (2021) aponta que com o crescimento acentuado da ciência de dados, da engenharia e da medicina, a integração entre ML e modelagens em multiescala podem apresentar soluções inovadoras para problemas complexos, sendo por isso cotada, junto com a computação de alta-performance, representantes de futuras áreas do século 21.

A modelagem de processos governados por leis da física através análises numéricas e simulações apresenta um caminho histórico que se inicia com equações diferenciais (HRENNIKOFF, 1941) e que hoje continua com os algoritmos paralelos de ML¹ (RADHAKRISHNAN, 2021), apresentando avanços significativos e hoje sendo aplicada em diversas áreas da química computacional, como: análises de reações químicas, otimização de propriedades e estruturas moleculares, descoberta de novos materiais, armazenamento de energia e predição de condições operacionais ótimas para processos químicos (MANN; VENKATASUBRAMANIAN, 2021).

Apesar das aplicações de algoritmos de ML na área de processos industriais ainda serem

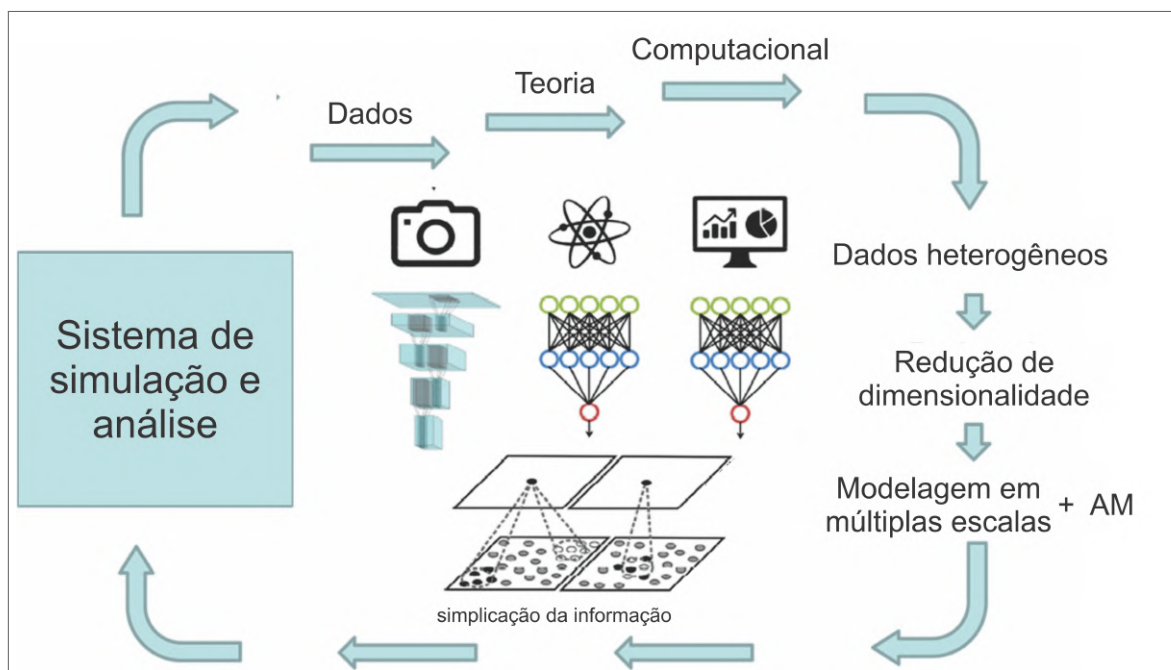
¹ <<https://cvw.cac.cornell.edu/APC/>>

pouco exploradas (LEE; SHIN; REALFF, 2018), o que se deve em particular aos modelos estatísticos de controle preditivos (MPC) bem estabelecidos por mais de 4 décadas, ferramentas de ML tem se mostrado extremamente úteis na explicação de desvios em modelos fundamentais da engenharia. Sendo esses desvios descobertos e compreendidos através da interpretação de comportamentos e padrões encobertos (COLEGROVE, 2020).

Por outro lado, a implementação de modelos de ML em processo industriais apenas por especialistas em AI podem levar a uma série de problemas práticos pela falta de conhecimento técnico em química, como a falsa necessidade de fluxo de dados e informações, a construção de modelos que não prevejam variações e limitações físicas e químicas que ocorrem nos processos industriais o que pode levar a ocorrência de erros na construção de modelos de aprendizado profundo (*Deep Learning*) (DL) (COLEGROVE, 2020).

Isso demonstra a necessidade de integração de conhecimentos técnicos e teóricos da engenharia e computacional para implementação de modelos de ML, para de forma contínua explorar o potencial desses modelos de acelerar, otimizar, e ajustar soluções para problemas multi-físicos conforme a Figura 15 (RADHAKRISHNAN, 2021).

Figura 15 – Construção de sistemas de modelagem em múltiplas escalas e modelos de aprendizado de máquina (ML) para acelerar criação de sistemas. A combinação desses dois fatores facilita a descoberta de modelos interpretáveis a partir de dados heterogêneos e irregulares, guiando a aquisição e interpretação física-química de novas informações.



Fonte: RADHAKRISHNAN (2021) adaptado.

2.3 IMPLEMENTAÇÃO AUTOMÁTICA DE MODELOS DE AM

Com o surgimento de modelos de ML cada vez mais complexos e densos em hiper-parâmetros, o processo de implementação e ajustes desses algoritmos se tornou uma atividade com consideráveis custos de tempo e desenvolvimento até mesmo para especialistas (HE; ZHAO; CHU, 2020). Por outro lado, o crescimento do poder computacional possibilitou implementações dinâmicas de diversos algoritmos de ML, e a resolução de problemas como otimização de hiper-parâmetros (*Hyperparameters Optimization*) (HPO) e a otimização de arquitetura (*Architecture Optimization*) (AO) (YANG et al., 2020).

Com o objetivo de facilitar a implementação dos algoritmos de ML para resolução desses problemas e torná-la uma tarefa mais acessível a outras áreas, foram criados modelos com aprendizado de máquina automático (*Automatic Machine Learning*) (AutoML), que propiciam a construção de modelos de ML sem a necessidade de manipulação profunda na arquitetura desses modelos (HE; ZHAO; CHU, 2020).

Modelos de AutoML servem tanto para a descrição de métodos automáticos de seleção de modelos e otimização de hiper-parâmetros quanto para a construção das etapas de preparação de dados, manipulação de variáveis, e geração e avaliação de modelos (HE; ZHAO; CHU, 2020; WENG, 2019). Recentes aplicações tem possibilitado a construção de arquiteturas de redes neurais, de modelos de aprendizado por reforço e também de algoritmos evolucionários (WENG, 2019), para serem usados na execução de tarefas como detecção de objetos e classificação de imagens.

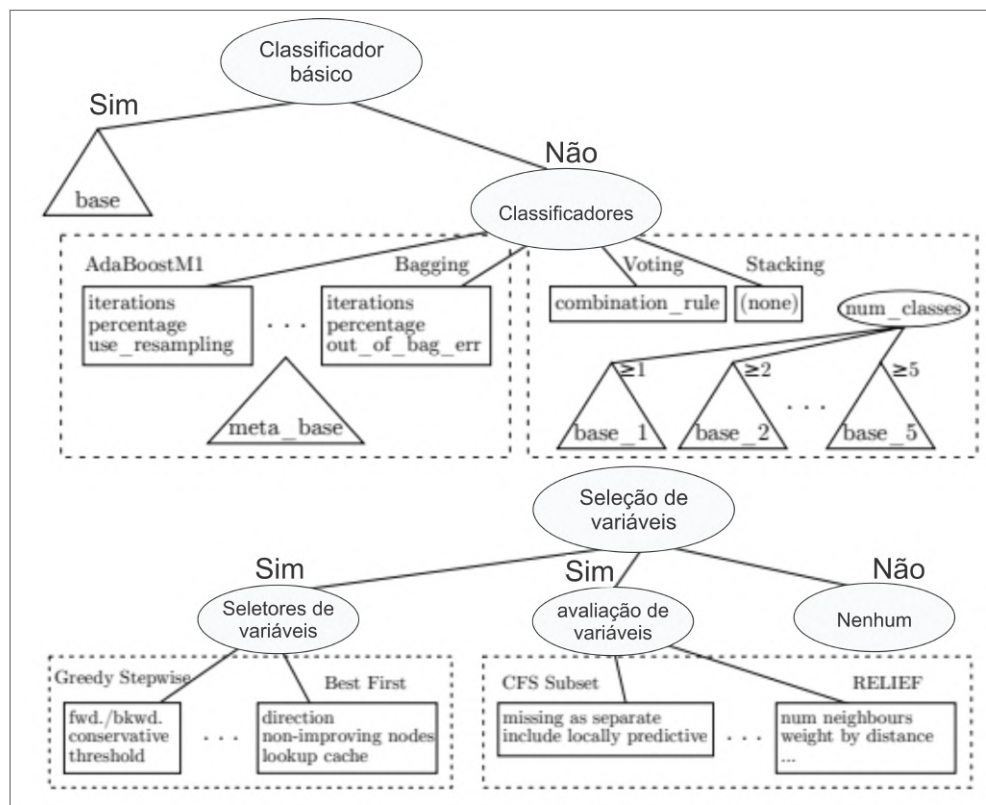
Dificuldades em relação ao alto consumo computacional e instabilidade desses algoritmos tem impulsionado as atuais pesquisas nessa área, como a de Chen et. al. (2019) com o desenvolvimento da pesquisa por arquitetura neural evolutiva reforçada (*Reinforced Evolutionary Neural Architecture Search*) (RENAS), e Cai et. al. (2018) com o uso de aprendizado por reforço como um meta-controlador para busca de arquiteturas a partir de modificações em redes e modificação de pesos. Outros trabalhos, como o de Liu et. al. (2018) tornam o espaço de busca de modelos contínuo, otimizando os seus modelos a partir dos valores obtidos pelo uso de gradientes descendentes, de forma mais eficiente que os modelos baseados em redes.

Como aplicação, Barreiro et al. (2018) propôs o *Net-Net* AutoML para predição de redes de ecossistemas biológicos, utilizando para isso informações de entropia de Shannon. Já Mohr et al. (2018) apresentou um novo modelo de AutoML baseado em planejamento hierárquico, que se mostrou competitivo frente a sistemas de AutoML de reconhecida eficiência, como Auto-

WEKA (THORNTON et al., 2012), Hyperopt-Sklearn (KOMER; BERGSTRA; ELIASMITH, 2019), TPOT (OLSON; MOORE, 2019) e Auto-Sklearn (FEURER et al., 2020).

O Auto-Weka utiliza de vários hiper-parâmetros para a construção dos seus modelos, os quais são escolhidos em intervalos de valores, contínuos ou discretos, diretamente correlacionados aos valores de precisão dos modelos. Além disso, seus modelos são otimizados a partir da configuração sequencial de algoritmos baseados em árvores, como os modelos gaussianos e os modelos de árvores aleatórias, correlacionando os seus valores de perda com os valores de hiper-parâmetros escolhidos (HUTTER; KOTTHOFF; VANSCHOREN, 2019). O Auto-Weka também permite a combinação de até 5 modelos de classificação, o que contribui significativamente para o tamanho total do espaço de configurações existentes para esse modelo de implementação automática (Figura 16).

Figura 16 – Espaço de configurações do Auto-Weka. Primeiramente é selecionado o método de classificação, a partir dos 27 modelos básicos, *base models*, de ML disponíveis (triângulos) ou por modelos com configurações específicas a serem definidas (elipses). Já a segunda etapa consiste na seleção das etapas de seleção e pre-processamento das variáveis.



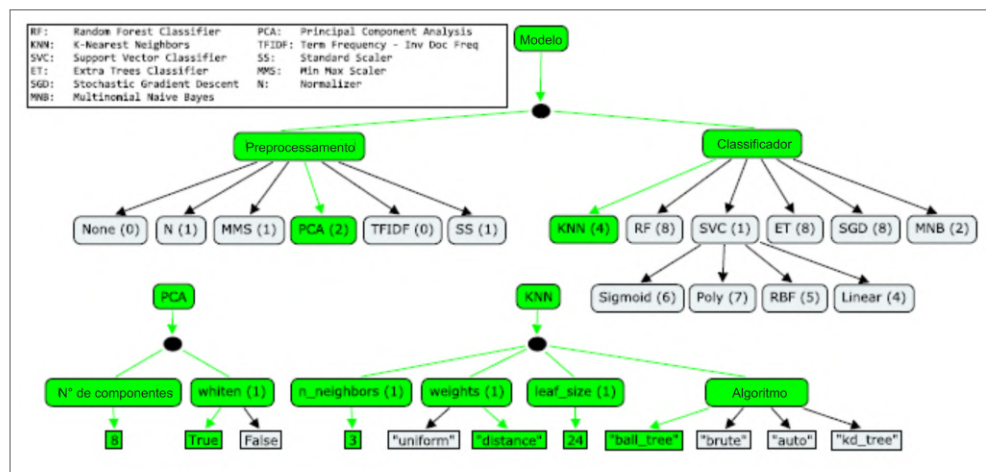
Fonte: THORNTON et al. (2012) adaptado.

Apesar do Auto-Weka ser um modelo bastante completo e que forneça a utilização de toda a biblioteca de modelos de aprendizado de máquina, ele não foi construído pensando em escalabilidade, criando a necessidade de exploração de modelos automáticos que suportem

aplicações científicas de forma escalável como o modelo Hyperopt-Sklearn (BERGSTRA et al., 2013).

O Hyperopt é construído nas linguagens Python e C, o que lhe confere maior velocidade. A sua biblioteca oferece algoritmos de otimização para a busca de configurações com diferentes tipos de variáveis, perfis de sensibilidade e estruturas condicionais, cuja aplicação exige a escolha de um domínio de busca, de uma função objetivo e de um algoritmo de otimização (Figura 17).

Figura 17 – Representação do espaço de busca do modelo automático Hyperopt-sklearn, que consiste na conjunção de aplicações disponíveis no sklearn de pré-processamento e de classificadores. Com 6 opções de cada, o hyperopt constrói suas opções de modelos a partir de amostragens aleatórias. Os blocos em verde mostram um exemplo de modelo construído, onde esse modelo (PCA e KNN) representam em conjunto o uso de 7 módulos do Hyperopt.

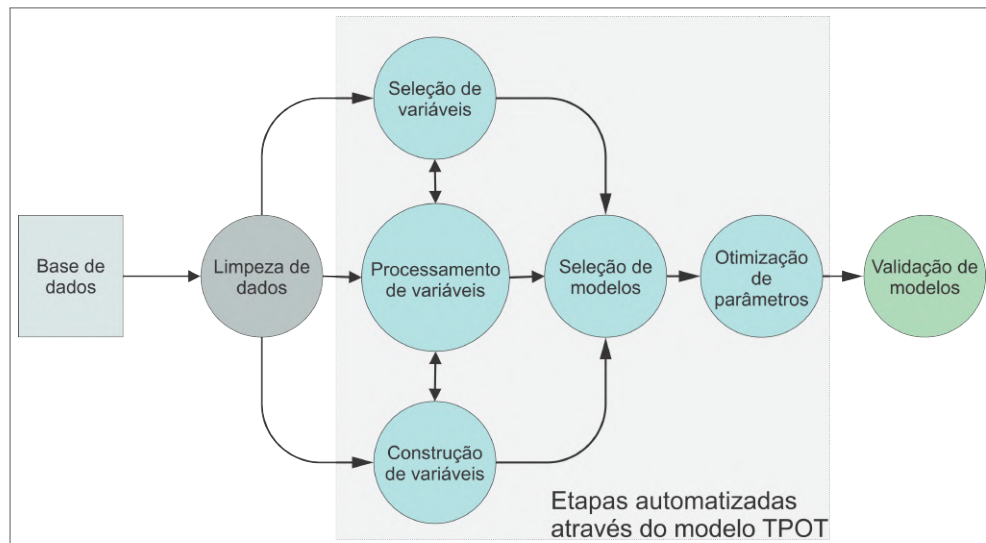


Fonte: BERGSTRA et al. (2013) adaptado.

Esse domínio de busca é especificado por variáveis escolhidas de forma aleatória, e as combinações mais promissoras são priorizadas a partir de maiores valores de probabilidade. A função objetivo por sua vez, associa amostras dessas variáveis aleatórias a valores de performance a serem otimizados (HUTTER; KOTTHOFF; VANSCHOREN, 2019).

Outra opção de sistema automático é a ferramenta de otimização com estrutura baseada em árvores (*Tree-Based Pipeline Optimization Tool*) (TPOT). Esse algoritmo otimiza a escolha de uma série de operadores para pré-processamento de variáveis e para a escolha de modelos de aprendizagem de máquina com o objetivo de maximizar a acurácia dos sistemas de classificação supervisionados. O TPOT trata a combinação dos operadores em um sistema de aprendizado de máquina como programas genéticos primitivos para associá-los em forma de árvores e otimizá-los por algoritmos genéticos. A Figura 18 mostra um exemplo do sistema de aplicação do TPOT, onde os círculos representam os seus operadores.

Figura 18 – Ilustração das etapas de construção de um sistema de ML através do modelo automático TPOT. Esse modelo automático busca a otimização das etapas de processamento de dados, de variáveis, da escolha dos modelos de ML e da otimização de hiper-parâmetros através de modelos baseados em estruturas de árvores.



Fonte: Elaborada pelo autor (2021).

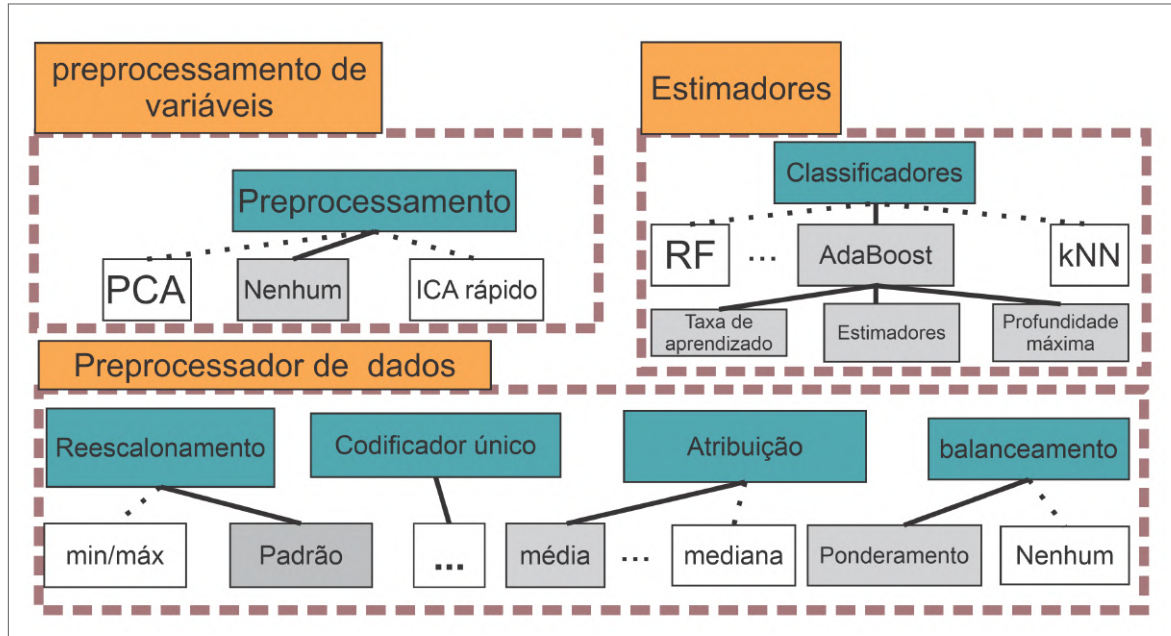
O que diferencia os modelos de AutoML e o seu processo automático de seleção combinada de algoritmos e de hiper-parâmetros (*Combined Algorithm Selection and Hyperparameter Optimization*) (CASH), é a capacidade de gerar previsões com restrições de tempo e memória. Em via de análise e comparação dos modelos de AutoML o desafio ChaLearn AutoML (GUYON et al., 2019) mostrou que o Auto-Sklearn (FEURER et al., 2020) apresentou os melhores desempenho, sendo por isso escolhido para este trabalho.

2.3.1 Auto-Sklearn

O Auto-Sklearn, assim como o Auto-WEKA, AutoKeras (JIN; SONG; HU, 2019) e Auto-PyTorch (ZIMMER; LINDAUER; HUTTER, 2020), utiliza da biblioteca *Scikit-learn* para composição dos seus sistemas (PEDREGOSA et al., 2011). Já o seu espaço de configurações possíveis é dado pela combinação entre as formas de pré-processamento de dados, de pré-processamento de variáveis e dos modelos de ML disponíveis.

As configurações apresentadas pelo Auto-Sklearn são organizadas hierarquicamente em forma de árvore contendo hiper-parâmetros contínuos, categóricos e condicionais (Figura 19). Com 15 modelos de classificação, 16 métodos de pré-processamento de variáveis e diversas aplicações de processamento de dados, totalizando cerca de 153 hiper-parâmetros (FEURER et al., 2020).

Figura 19 – Espaço de configurações do modelo de implementação automática Auto-Sklearn. Através da análise de performance e da otimização bayesiana de hiper-parâmetros, são construídos os sistemas de predição, dados pela combinação das etapas de processamento de dados, variáveis e pela escolha do modelo de ML. Em azul são os hiper-parâmetros principais, já em cinza são os hiper-parâmetros secundários escolhidos para exemplo de possível configuração para um sistema de ML.



Fonte: ONO et al. (2020) adaptado.

Os algoritmos de classificação podem ser separados em diferentes categorias, como modelos lineares, SVM, análise de discriminantes, KNN, árvores de decisão, e os modelos combinados, também conhecidos como *ensembles* (FEURER et al., 2015). E diferentemente de outros modelos automáticos, como o Auto-Weka (THORNTON et al., 2012), o Auto-sklearn apresenta prioridade para o uso de classificadores básicos para treinamento e a análise de combinação com modelos treinados, o que se mostrou mais eficiente (FEURER et al., 2015).

Já as etapas de preprocessamento são divididas em duas partes, preprocessamento de dados e preprocessamento de variáveis. O preprocessamento de dados compreende o reescalonamento das entradas, o preenchimento de dados faltantes, transformação dos rótulos em valores binários e o balanceamento de classes. E por último, os tipos de preprocessamentos de variáveis disponíveis são categorizados em seleção de variáveis, aproximadores para os núcleos de processadores, decomposição matricial, acondicionamentos, agrupamento de variáveis, expansões polinomiais de variáveis e métodos de seleção de variáveis (FEURER et al., 2015). Tendo assim à disposição um espaço de configurações formado pelo modelo automático de implementação de combinação de estimadores, dos preprocessamentos das variáveis e dos preprocessamentos de dados, e pelos hiper-parâmetros de tempo de estimação e de tempo de treinamento dos

sistemas de predição.

O mecanismo de construção de modelos no Auto-Sklearn é, principalmente, baseado na avaliação de tempo, sendo a busca por configurações guiada pela otimização Bayesiana (*Bayesian Optimization*) (BO) (SHAHRIARI et al., 2016). A BO consiste na avaliação de iteração de duas etapas: 1) Correlação entre modelos probabilísticos, e seus hiper-parâmetros, com resultados de perda. 2) Otimização na escolha de funções de busca de novas configurações de hiper-parametros para avaliação. Diferentemente da maioria dos AutoML que utilizam modelos Gaussianos como preferência (RASMUSSEN; WILLIAMS, 2005), o Auto-Sklearn busca a utilização de modelos com combinações de árvores de decisão, devido a melhores performances em problemas com alta dimensionalidade e otimização estruturada, característico de problemas do tipo CASH (HUTTER; HOOS; LEYTON-BROWN, 2011).

2.4 VISUALIZAÇÃO DE VARIÁVEIS E MODELOS DE ML

O crescimento acelerado de aplicações de modelos de ML na indústria é dado pela capacidade desses modelos de extrair padrões e formar predições em cenários complexos. Por outro lado, os modelos de ML com maiores capacidades de modelagem de sistemas também apresentam alta complexidade em sua construção, levando a necessidade de ferramentas que auxiliem no entendimento e na extração de informações e valores, como as técnicas de visualização (ZHOU et al., 2018).

2.4.1 Métodos de interpretação de modelos de ML e variáveis

Parte importante dos modelos avançados de ML são os baseados em aprendizado profundo, *deep-learning*, também conhecidos como modelos do tipo caixa-preta, devido a sua forma combinatória ser de difícil interpretação (KAHNG et al., 2017). Considerando isso, diversas ferramentas de visualização tem buscado facilitar a análise desses modelos, como o ViDx (XU et al., 2017), ActiVis (KAHNG et al., 2018) e Squares (REN et al., 2017), fornecendo avanços no processo de interpretação de modelos de aprendizado profundo.

Com o avanço das técnicas de visualização para suporte a análise de modelos, surgiram visualizações capazes de auxiliar a interpretação de modelos automáticos de ML, os AutoML, como o de Ono et al. (2020) que permite a observação de combinações dos métodos de pré-processamento, extração de variáveis, operadores e de classificadores. Com esse tipo de

visualização, é possível além de implementar modelos de ML de forma automática, observar diferentes configurações para predição e compará-las através de métricas de performance.

Tão importante quanto interpretar os modelos de ML, é entender as variáveis e relevâncias para os modelos de predição, pois, além de construir as correlações observadas nos modelos de ML, são esses parâmetros que regem o resultado final obtido pelo modelo, o que em processos industriais representa a qualidade final do produto obtido no processo industrial. Ferramentas de visualização como explanação aditiva de Shapley (*SHapley Additive Explanations*) (SHAP) (ZHANG; ZHANG; WANG, 2012), permitem a análise da importância das variáveis para a construção da predição do modelo, assim como, o impacto sobre a predição de acordo com os seus valores.

Os valores de SHAP, θ , para as variáveis j são calculados de acordo com a observação média dos valores de importância ou impacto dessas variáveis. Que são determinados pela diferença entre os valores de predição f com a variável j inclusa na função de predição e os valores de predição de f sem que haja a variável j na função (Equação 2.3) (LUNDBERG; LEE, 2017).

$$\begin{aligned}\theta_j^m &= f(x_{+j}^m) - f(x_{-j}^m) \\ \theta_j(x) &= \frac{1}{M} \sum_{m=1}^M \theta_j^m\end{aligned}\tag{2.3}$$

Esse cálculo é feito de forma repetida para cada instância x em um conjunto menor de amostras z pertencente a base de dados.

Já as bibliotecas Plotly (INC., 2015) e Skater (CHOUDHARY; KRAMER; TEAM, 2018), permitem a análise de objetos que compõem o problema e as suas correlações, para observação de padrões, valores de variáveis e valores de impacto, SHAP, que favoreçam uma classe em detrimento da outra.

O uso de diferentes bibliotecas de visualização para este trabalho se deve ao fato que cada tipo de visualização apresenta, intrinsecamente, um propósito principal (MIDWAY, 2020). As visualizações como os gráficos de dispersão, gráficos de barras, tabelas e plotes paralelos interativos, buscaram, respectivamente, ressaltar correlações, importâncias, informações e padrões presentes nos dados.

2.4.2 Métodos de visualização de objetos

O uso de monitoramento de processos, como detecção, diagnóstico e prognóstico de falhas, em tempo real é crucial para a garantia de qualidade e de produção na indústrias (YUE et al., 2000). Para essa atividade, diversos métodos de visão de máquina tem sido aplicados na indústria, como o ExtremeNet (ZHOU; ZHUO; KRÄHENBÜHL, 2019) utilizado na detecção de instâncias, e o YOLOv3 (REDMON; FARHADI, 2018) que utiliza de retângulos para detecção de bordas e enquadramento de objetos. Aplicações como essas podem ser relevantes para a identificação e marcação de contornos e objetos, utilizados em geral para cenários com informações simples, e alta discriminação (LEI et al., 2020).

Em particular, para o processo industrial de recobrimento de filme por deposição metálica a vácuo os valores de densidade óptica (*Optical Density*) (OD), forma de medida indireta da espessura de cobertura metálica sobre o filme polimérico, são continuamente obtidos em tempo de processo e o fato do recobrimento ser feito através de evaporadores em paralelo, assim como os sensores de OD, faz com que essas medições físicas representem a qualidade dos sistemas de deposição de forma local. Isso torna a posição espacial de objetos um termo importante para análise da qualidade do sistema de recobrimento local e para detecção de possíveis defeitos de recobrimento, dados por erros no sistema ou pela presença de furos, sujeiras e arranhões (BISHOP, 2011). Assim, é intuitivo observar do perfil de valores de OD obtidos em cada objeto como forma de análise de qualidade do recobrimento, utilizando para isso técnicas de visão computacional para compreensão de imagens e observação de objetos, como a segmentação semântica (*Semantic Segmentation*) (SS) (WANG; LIU; XU, 2017).

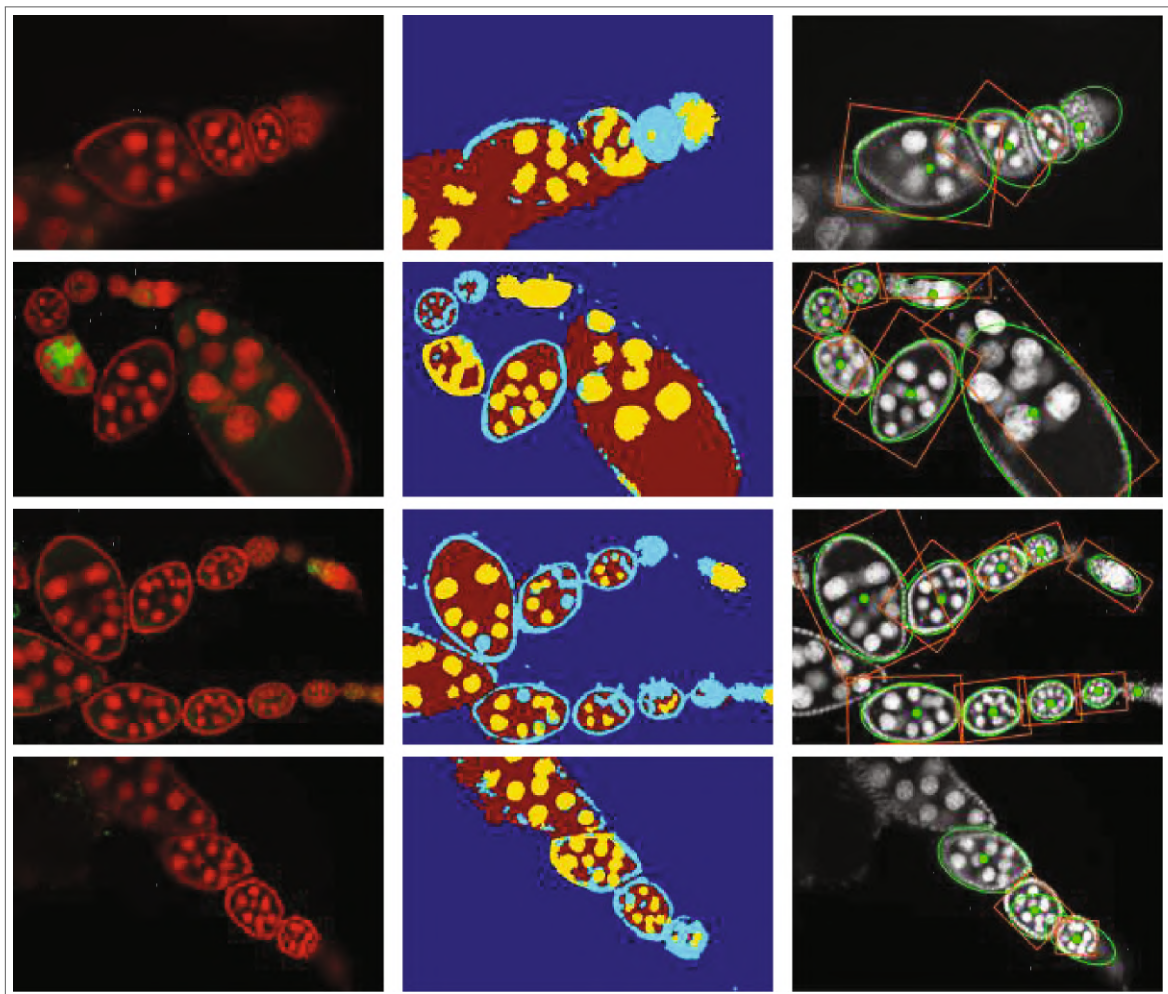
2.4.2.1 Segmentação semântica

Entender cenários e imagens não é uma atividade simples, principalmente, quando existem informações variadas e de difícil distinção. E a isso as ferramentas de SS se propõem solucionar (EIGEN; FERGUS, 2015), utilizando-se de informações locais da própria imagem para segmentação e informações globais para classificação (LONG; SHELHAMER; DARRELL, 2015). A SS de imagens pode ser feita de diversas formas e por diferentes algoritmos, modelos do tipo DL para detecção e classificação de objetos são bastante conhecidos e amplamente utilizados (GIRSHICK et al., 2016). Por outro lado, modelos de SS baseados em super-píxeis associados a modelos tradicionais de ML tem crescido em aplicação, e inclusive apresentado desempenhos

ótimos em competições como a PASCAL VOC Challenge (FULKERSON; VEDALDI; SOATTO, 2009; GOULD et al., 2008; YANG et al., 2010) para a estimativa de profundidades (ZITNICK; KANG, 2008), e para a segmentação (LI et al., 2004) e localização de objetos (FULKERSON; VEDALDI; SOATTO, 2009).

Borovec et. al. (2017) utilizou a segmentação baseada em super-píxeis para computação de cor e textura em imagens de infravermelho de ovos de *drosophila*, pequenas moscas que sobrevoam frutas, atribuindo aos super-píxeis das imagens cores de acordo com as classes que lhe foram atribuídas por um classificador baseado na combinação de árvores de decisão, seção 2.2.1.1, com regularização por Graphcut (BOYKOV; VEKSLER; ZABIH, 2001) (Figura 20).

Figura 20 – Exemplo de segmentação semântica em quatro classes com divisão de imagem em super-píxeis e agrupamento por classes através de um classificador de árvores aleatórias, utilizando de imagens de infravermelho de ovos de *Drosophila* obtidas em microscópio eletrônico. A primeira coluna de imagens mostra as imagens originais, a segunda a segmentação em classes representadas por diferentes cores e a terceira coluna demonstra a construção supervisionada dos seguimentos e a localização de objetos presentes na imagem.



Fonte: BOROVEC; KYBIC; NAVA (2017).

É possível observar na Figura 20 a distinção por classes das regiões das imagens de infra-

vermelho, além da construção de contorno de objetos, algo que seria de difícil diferenciação entre os pixels das imagens originais.

Com esse objetivo, a ferramenta utilizada por Borovec et al. (2017) foi implementada para análise dos perfis de OD dos produtos obtidos durante o processo de metalização do filme. Essas imagens foram formadas a partir dos dados obtidos pelos sensores de OD distribuídos transversalmente ao filme e captados ao longo de todo o comprimento dos filmes a medida que eles estavam sendo recobertos, formando assim, conjuntos matriciais de informações de OD que foram convertidos em imagens.

Da mesma forma que as imagens de infravermelho dos ovos de *drosophila*, as imagens originais formadas com as medidas de OD apresentam dificuldades para diferenciação dos pixels. E a segmentação semântica possibilita uma melhor interpretação desses produtos, através da formação de contornos e da segmentação da imagem em classes para serem analisadas por especialistas. Além disso, a formação de segmentos na imagem contribuem para o melhor treinamento de modelos complexos de ML, que podem ser utilizados para predição de variáveis dos produtos formados.

3 METODOLOGIA

3.1 COLETA DE DADOS

Como mostrado na seção 2.1.3, os dados coletados para esse projeto consistem na combinação de informações de produção, das condições operacionais e de medições de propriedades do produto em laboratório, como aspectos de qualidade dos filmes produzidos e a sua classificação final. Sendo referentes a produção obtida no período entre setembro de 2019 e agosto de 2020.

Os dados de qualidade dos produtos formados a partir de análises laboratoriais apresentam altos custos de tempo para sua obtenção, além de serem feitos em pequenos recortes do filme produzido, o que pode levar a tempo de espera elevado e a obtenção de informações não representativas devido a variabilidade espacial apresentada pelo recobrimento. Assim, partindo do entendimento de reduzir o tempo para obtenção de informações de forma a agilizar o processo de monitoramento da qualidade, foi escolhida apenas a variável de classificação final do produto dentre as obtidas através de análises laboratoriais.

Já as informações de produção serviram como forma de cruzamento de informações, pois, apresentavam os códigos da amostras que eram enviadas para análises laboratoriais e também o registro de produção do equipamento que é armazenado no disco-rígido da máquina.

Utilizando o SQL Server Studio® foi possível extrair e reestruturar os bancos de dados armazenados no equipamento industrial. O primeiro, com os parâmetros de processo como temperatura do tambor de resfriamento, pressão interna da zona de recobrimento e entre outras variáveis com valores médios presentes no Quadro 1 da subseção 2.1.4. E o segundo banco de dados com as informações de OD dos filmes produzidos, com intervalo de valores de 2 min e o comprimento referente aos dados captados. Já os dados de produção e de laboratório são armazenados no programa TOTVS®, de onde foram extraídas em forma de tabelas. Para ambos os bancos de dados, foram retirados objetos repetidos ou com valores nulos.

3.2 PREPROCESSAMENTO DOS DADOS

Através da coleta de dados em equipamentos e sistemas, foram feitos cruzamentos de informações de produção, de processo e de qualidade. Ao final, foram obtidas duas tabelas de dados, de onde foram excluídos exemplos duplicados e exemplos com variáveis nulas. Para os

valores das variáveis de processo, foram observados os valores que destoavam das médias das variáveis, *outliers*, com valores absolutos maiores que três desvios-padrões em relação a média, o que demonstrou que a sua exclusão reduzia a capacidade de generalização dos modelos, fornecendo acurácia de 0.847 para o modelo de maior performance, o que levou a optar pela manutenção desses valores. O primeiro conjunto de variáveis são referentes a informações de condições de processo com as respectivas classificações dos exemplos. E o segundo conjunto de dados descrevem os valores de OD do recobrimento dos produtos obtidos com os respectivos valores de comprimento, e também com a classificação final desses produtos.

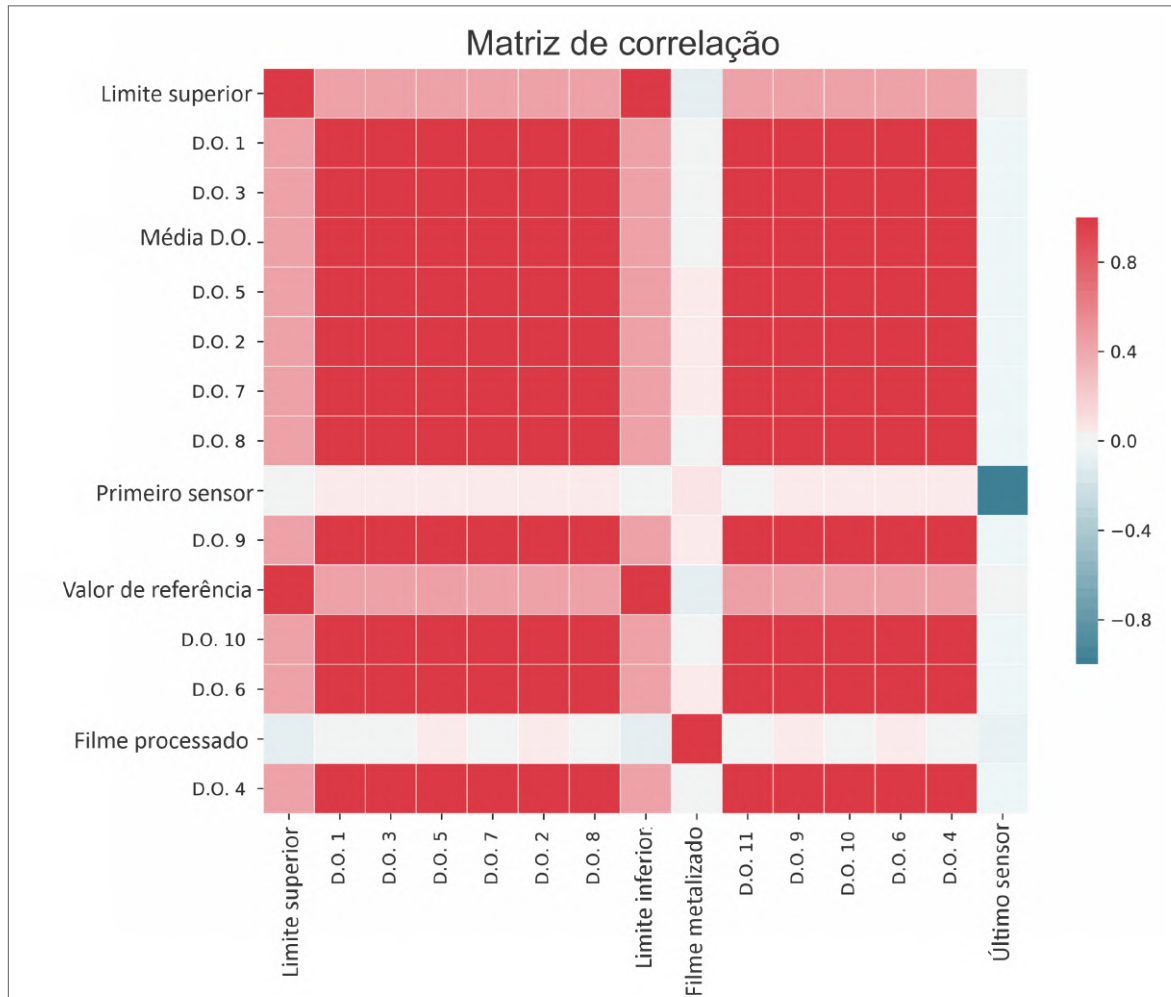
3.2.1 Seleção de variáveis

A Tabela 1 é referente às informações de condições operacionais para os produtos obtidos no processo de metalização à vácuo. Diferentemente da Tabela 2, a tabela de dados de processo apresenta alta dimensionalidade, e como já citado anteriormente, a reduzida quantidade de objetos com classificação (rótulo) fazem com que os modelos de ML apresentem dificuldades para construção do modelo. Assim, foram retiradas variáveis de identificação e variáveis categóricas, permanecendo no banco de dados as informações de processo e informações médias de OD. Como forma de melhorar o processo de treinamento dos algoritmos na Tabela 1, foram observadas as variáveis que apresentavam alta correlação e baixa contribuição de informação para o treinamento dos modelos de ML. A Figura 21, apresenta as variáveis do processo processo que apresentam correlação maior de 97,5%.

É possível observar que parte representativa das variáveis que apresentam correlações altas são referentes aos valores médios de OD, isso pode ser compreendido tendo em vista que essas variáveis tem o objetivo de atingir valores de recobrimento iguais de forma a manter o perfil do filme uniforme. Uma vez observada a existência de variáveis com correlações elevadas, foi feita uma redução de dimensionalidade da base de dados através da ferramenta *feature-selector* <<https://github.com/WillKoehrsen/feature-selector.git>> excluindo variáveis de baixa relevância e de forma a manter 99% da representatividade do conjunto de dados. A Tabela 1 mostra a descrição das 24 variáveis restantes de um conjunto original de 41 variáveis.

Os banco de dados com as variáveis de processo constam com 494 objetos dos quais cerca de 85% são pertencentes a classe de filmes aprovados, situação igual a 1, porém, para esse banco de dados as etapas de balanceamento e outros preprocessamentos de dados e variáveis foram aplicadas de forma automática por meio do Auto-sklearn (seção 4.1). A depender do

Figura 21 – Matriz de correlação com variáveis de processo com correlação maior que 97,5%. A presença de variáveis com correlações altas é dada, principalmente, pelos valores de densidade óptica obtidos nas diferentes posições dos sensores. Isso pode ser entendido, uma vez que esses valores de recobrimento variam em torno de um mesmo valor pré-determinado, fazendo com que os valores de OD acabem apresentando variações coordenadas, mesmo com os sistemas locais de recobrimento agindo de forma independente.



Fonte: Elaborada pelo autor (2021)

modelo utilizado, o desbalanceamento dos dados pode levar a falhas de predição, como a escolha indiscriminada pela classe majoritária, uma vez que devido a sua escolha em relação às demais classes, leva a obtenção de taxas altas de desempenho.

A Figura 22 mostra a dispersão dos valores de OD que variam de 0 a 5, referentes as posições das barcas de evaporação B2, B3, B4, B23, B24 e B25. É possível observar que os valores de recobrimento válidos para pouco mais de 75% dos objetos são restritos as posições de B3 a B24, o que fez optar pela escolha da exclusão das variáveis B2 e B25.

Já na Tabela 2 é possível observar que os comprimentos dos objetos são variáveis, com média de 19 km e com objetos que chegam até 87 km de extensão, o que nos fez optar por

Tabela 1 – Descrição do banco de dados formado a partir das variáveis do processo industrial de metalização de filme por deposição. As colunas apresentam os valores médios e os valores presentes nos quantis. O que permite a observação das condições operacionais utilizadas na obtenção dos filmes metalizados. A variável situação é o rótulo correspondente a classificação final dos produtos, sendo o valor um para produtos aprovados e zero para produtos reprovados.

	média	min	25%	75%	max
Furos P	1.45	0.10	0.69	1.98	6.36
Furos M	0.39	0.02	0.24	0.49	1.90
Furos G	0.22	0.00	0.12	0.28	1.14
Arranhão	0.04	0.00	0.01	0.05	0.72
Desvio Padrão DO ^a	0.02	0.01	0.01	0.02	0.05
Média DO	3.02	2.39	2.42	3.42	3.46
Tempo total (min)	187.04	87.93	160.33	195.72	875.43
T ^b de aquecimento (min)	11.78	3.70	8.30	14.72	27.33
T. de recobrimento (min)	100.16	49.17	94.84	107.03	131.33
Filme inspecionado (m)	30505.28	0.00	29345.25	36064.50	44621.00
T. de resfriamento (min)	124.82	66.90	120.78	132.72	192.40
T. das bombas (min)	12.01	7.57	9.81	12.98	26.27
Tempo de ventilação (min)	6.05	2.97	3.77	9.23	14.87
Filme processado (m)	36561.50	17173.00	36428.00	36813.00	45010.00
T. de máquina aberta (min)	48.71	2.67	24.43	54.48	731.50
T. de vida das barcas (min)	81.45	1.65	20.91	113.98	359.10
Primeiro Ponta	3.29	1.00	3.00	4.00	6.00
Peso do rolo (kg)	1238.30	858.00	1159.00	1323.25	1515.00
Largura do filme (m)	2012.63	1400.00	1892.50	2135.00	2450.00
Número alvo	3.08	2.40	2.40	3.40	3.40
Vel ^c do filme (m/min)	365.16	300.00	350.00	380.00	530.00
Vel. do fio (m/min)	69.15	43.50	63.20	74.00	84.80
Tensão aplicada (V)	416.07	234.00	360.00	462.75	622.00
Tensão medida (V)	274.04	137.00	238.00	305.00	412.00
Situação	0.85	0.00	1.00	1.00	1.00

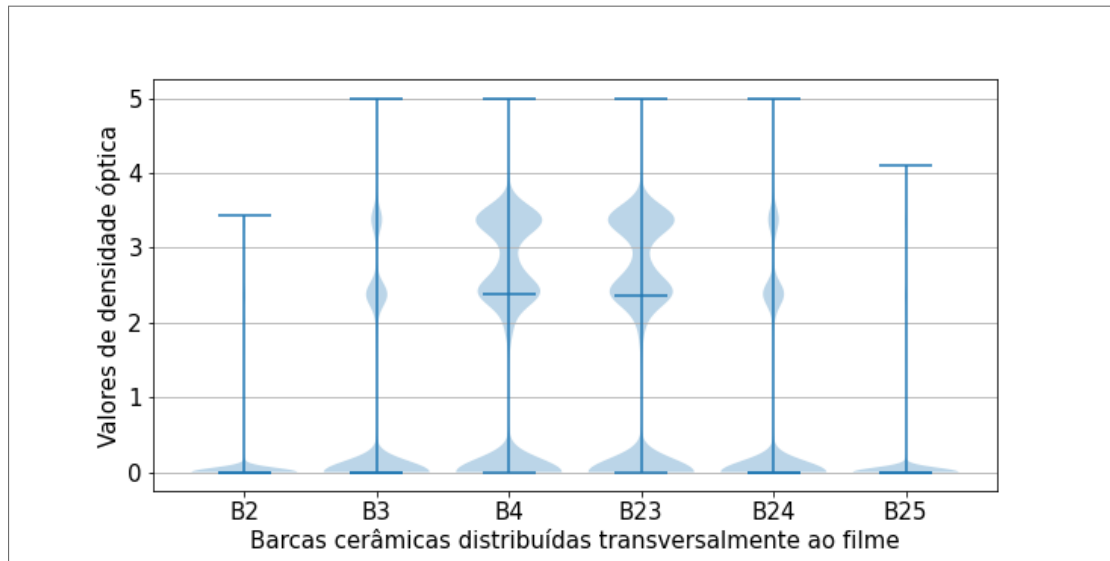
^aDensidade Óptica, ^bTempo, ^cVelocidade,

Fonte: Elaborada pelo autor (2021)

um comprimento máximo de objeto de 45 km, suficiente para compreender grande parte dos comprimento dos objetos mais que 75%, com adição de linhas nulas para objetos menores e com corte de informação para objetos maiores que 45 km, sendo por fim utilizados objetos com tamanho de 22 colunas por 65 linhas.

Como os filmes apresentam comprimentos e larguras variadas, as dimensões de imagens de

Figura 22 – Distribuição dos valores de densidade óptica dos objetos. Sendo cada coluna referente a posição do sistema de deposição onde há o recobrimento metálico e também as medidas de OD ao longo do comprimento do filme.



Fonte: Elaborada pelo autor (2021)

22x65 construídas foram escolhidas de forma a abranger as dimensões comumente utilizadas para os filmes produzidos pela TOPMET 2450, que são de até 4,5 m de largura e 47 km de comprimento.

Tabela 2 – Descrição da base de dados formada com valores de densidade óptica (*Optical Density*) (OD). Os valores foram obtidos a partir de 26 sensores ópticos distribuídos transversalmente ao comprimento do filme, com registro de informação a cada dois minutos e variação de zero a cinco, o que representa o menor e o maior valor de recobrimento, respectivamente. Associado aos valores de OD também são obtidos os respectivos valores de comprimento das medições.

	Comprimento (m)	B2	B3	B4	B23	B24	B25
Contagem	32240	32240	32240	32240	32240	32240	32240
Média	19386	0.09	0.68	1.71	1.72	0.69	0.10
Mínimo	0.00	0.00	0.00	0.00	0.00	0.00	0.00
25%	9650	0.00	0.00	0.00	0.00	0.00	0.00
50%	19131	0.00	0.00	2.38	2.37	0.00	0.00
75%	28733	0.00	1.91	3.25	3.25	1.99	0.00
Máximo	87259	3.43	5.00	5.00	5.00	4.99	4.11

Fonte: Elaborada pelo autor (2021)

As 32.240 linhas de informação, são referentes a 142 produtos, os quais foram utilizados em sua totalidade para a segmentação semântica. Dentre esses produtos analisados, 21 pertencem a classe de filmes desaprovados. Como forma de balancear a quantidade de objetos por classe,

foi feita uma sobre-amostragem na classe minoritária, equalizando a quantidade de objetos por classe. A pouca quantidade de objetos para treinamento dos modelos de classificação se deve a limitação de amostras que são escolhidas para análise laboratorial.

3.3 TREINAMENTO DOS ALGORITMOS

Os algoritmos de ML neste trabalho estão envolvidos em 3 etapas, a primeira compreende a aplicação de classificadores presentes na biblioteca `scikit-learn`¹ através do AutoML para exploração de modelos de classificação de produtos utilizando os dados de processo presentes na Tabela 1. Desses dados, foram separados aleatoriamente 75% para treinamento e 25% para validação.

A segunda etapa foi a construção da segmentação semântica do perfil de OD através do modelo de modelo de mistura de gaussianas (*Gaussian Mixture Model*) (GMM)(REYNOLDS, 2009). O agrupamento dos valores de OD foi feito em quatro classes e de forma não-supervisionada, segmentando as imagens de acordo com a co-variância local dos dados de OD.

E a última etapa de exploração dos modelos de ML busca classificar os objetos a partir dos perfis de OD segmentados e não-segmentados, utilizando para isso um modelo sequencial de redes neurais². Esse modelo de ML foi escolhido devido a sua capacidade de adaptação em problemas de alta dimensionalidade e a possibilidade de interpretação do modelo através da observação dos pesos presentes nas suas camadas internas. Para o treinamento do modelo de rede neural foram utilizados 80% dos exemplos disponíveis e selecionados de maneira aleatória, já os 20% restantes foram utilizados para validação do modelo.

Todos os códigos desenvolvidos foram implementados através da plataforma do Google Colab[®], e estão disponíveis no repositório *Master-project-code*³ no Github.

3.4 VISUALIZAÇÃO DOS RESULTADOS

Os modelos de visualização propostos compreendem os objetivos de prover a interpretabilidade dos modelos de ML construídos e de analisar a importância das variáveis e suas correlações com os parâmetros de qualidade. O primeiro modelo de visualização foi o Pipeline Profiler (ONO et al., 2020), utilizado para análise dos sistemas formados pela combinação entre

¹ https://scikit-learn.org/stable/supervised_learning.html#supervised-learning

² https://keras.io/guides/sequential_model/

³ <https://github.com/tmrbr/Master-project-codes>

as etapas de préprocessamento de dados e variáveis, e pela etapa de seleção de classificador feita pelo AutoML. De forma complementar, foram utilizadas as visualizações de gráficos paralelos e de distribuição de probabilidade disponíveis na biblioteca matplotlib para análise de performance dos sistemas mais bem classificados ⁴.

A etapa de visualização compreende a análise de importância das variáveis para o processo de predição de classe, utilizando para isso os valores de SHAP e as aplicações disponíveis na sua biblioteca⁵. Então, foram gerados dois gráficos com esses valores, um gráfico de barras horizontais e um gráfico com dispersão unidimensional das variáveis que compõem o processo de produção. E ainda com os valores de SHAP, foi feita a análise do perfil de OD e do modelo de rede neural utilizado, observando a correlação entre a distribuição dos valores de SHAP e a ocorrência de variações de OD no produto formado.

Ainda sobre as variáveis de processo, foram utilizadas as bibliotecas Plotly⁶ e Skater⁷ para visualização de objetos, análise de correlações das variáveis sobre a variável de maior impacto na predição de classe do modelo de melhor desempenho, e a visualização sistemática de variáveis e impactos sobre a classificação de um objeto desejado, respectivamente.

Já para visualização dos perfis de OD dos filmes formados, foram utilizadas a segmentação semântica (BOROVEC et al., 2017), a contagem de classes e a observação do modelo de redes neurais aplicado, tanto pela visualização de pesos ou valores de Shap (LUNDBERG; LEE, 2017) quanto pelas saídas das múltiplas camadas do modelo através da biblioteca Matplotlib.

As visualizações implementadas buscam complementar as análises de performances dos modelos de ML explorados nesse projeto, tanto para os modelos construídos a partir das variáveis de processo, quanto para o modelo de rede neural utilizado para classificação dos produtos a partir dos dados de OD. O que torna a análise de performance mais robusta, para além da análise das métricas de desempenho AUC, F1, acurácia, Matthews, precisão e perda de Brier, uma vez que permitem a comparação do impacto das variáveis e suas predições, com a real importância das variáveis do processo e a qualidade do produto avaliado.

⁴ <https://matplotlib.org/>

⁵ <https://shap.readthedocs.io/en/latest/index.html>

⁶ <https://plotly.com/python/>

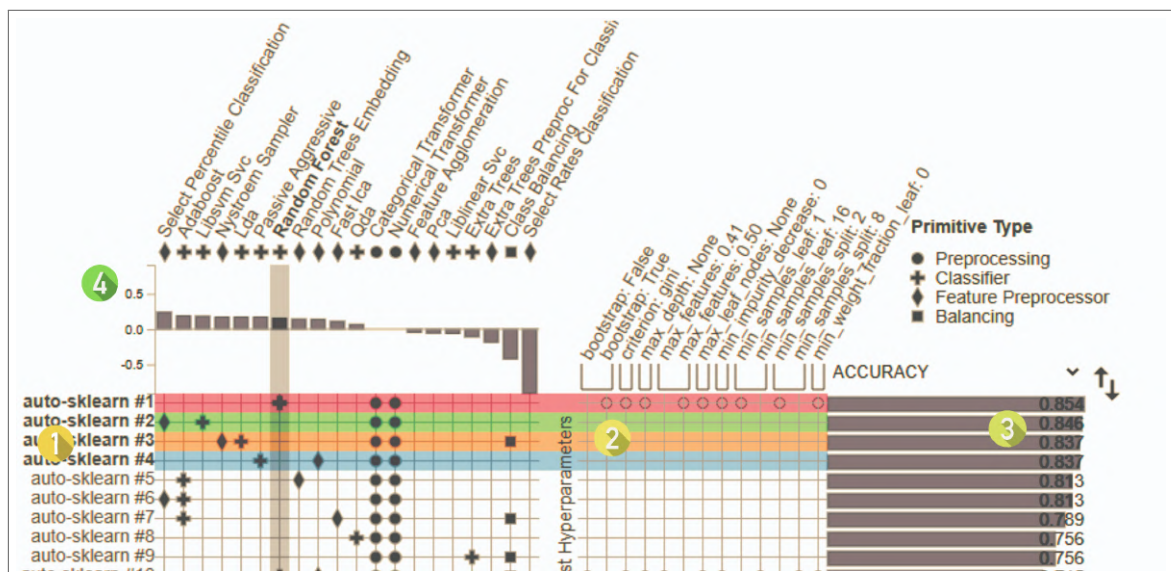
⁷ <https://oracle.github.io/Skater/index.html>

4 RESULTADOS E DISCUSSÃO

4.1 AUTOML

Para exploração dos modelos de ML utilizando o AutoML, os 494 objetos com as variáveis de processo foram divididos de forma aleatória e estratificada em 75% para exemplos de treinamento e 25% para teste e avaliação de performance dos sistemas formados por validação cruzada. O Auto-sklearn foi aplicado de forma padrão com tempo de atividade de 300 s e limite de 30 s por treinamento de modelo, de forma a obter modelos robustos e com baixo gasto computacional. A título de análise, tempos maiores de aplicação com até 5 horas de atividade e 600 segundos de treinamento, mostraram que há formação de modelos mais complexos e com maiores tempos de treinamento, porém as performances não melhoram para além dos modelos obtidos com os tempos estabelecidos como padrão. Assim, os sistemas treinados em tempo padrão estão dispostos abaixo com os respectivos valores de acurácia (Figura 23).

Figura 23 – Pipeline Profiler aplicado para visualização de sistemas de predição de classe dos filmes obtidos pelo Auto-sklearn aplicado a variáveis do processo de metalização a vácuo. A visualização é composta por: 1) Sistemas do Auto-sklearn formados pela combinação dos modelos primitivos de processamento de dados, classificadores e balanceamento, descritos pelas colunas e símbolos e ordenados de acordo com a sua performance. 2) Descrição dos hiper-parâmetros utilizados para o sistema selecionado. 3) Valores de performance de acordo com os sistemas construídos. Sendo também possível observar os modelos ordenados pelo tempo de predição. 4) Contribuição de cada etapa para a performance do sistema.



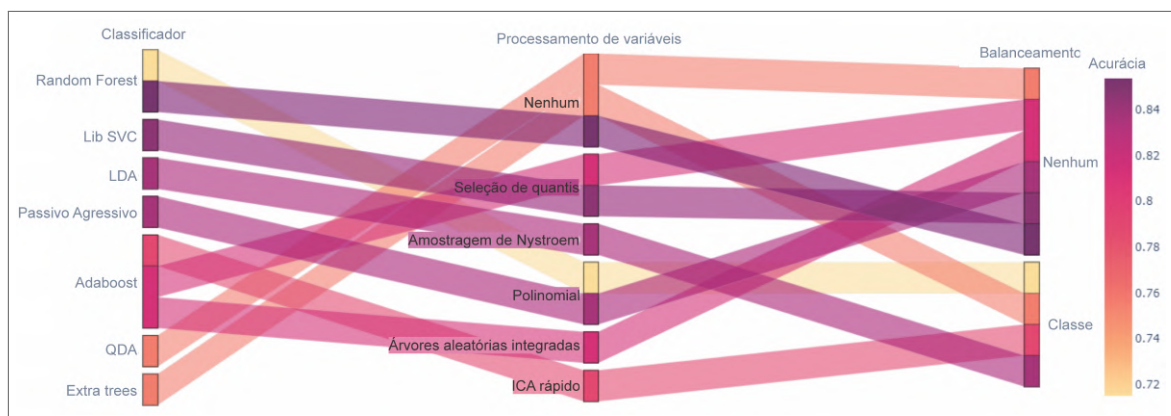
Fonte: Elaborada pelo autor (2021)

A nível de exploração, repetindo o processo de treinamento e avaliação, os resultados obtidos com maior frequência apresentaram o classificador florestas aleatórias (*Random Forest*)

(RF) com melhor performance, apresentando acurácia de 0.854 apenas com pré-processamento de dados e sem necessidade de balanceamento de classes. Seguido do sistema com classificador de vetores de suporte (*Support Vector Classifier*) (SVC) com acurácia de 0.846 e do análise por discriminante linear (*Linear Discriminant Analysis*) (LDA) com 0.837. É possível observar que a alta dimensionalidade do problema favorece um melhor desempenho de modelos lineares e que apesar de poucos exemplos de treinamento os modelos formados pelo Auto-sklearn apresentaram bons desempenhos em relação a métrica acurácia.

É possível observar que para os sistemas mais bem posicionados, foram feitas transformações numéricas e categóricas para a correta manipulação de variáveis de acordo com a necessidade do modelo de ML. Já para as etapas de processamento de dados e de processamento de variáveis foram utilizadas diferentes metodologias, impactando no desempenho dos sistemas de acordo com o mostrado na Figura 24.

Figura 24 – Representação dos dez sistemas com maior desempenho fornecidos pelo Auto-sklearn. As configurações dos sistemas de maiores desempenhos demonstram o uso de diferentes combinações de modelo de ML, pré-processamento de variáveis e de pré-processamento de dados, com os respectivos valores de desempenho.



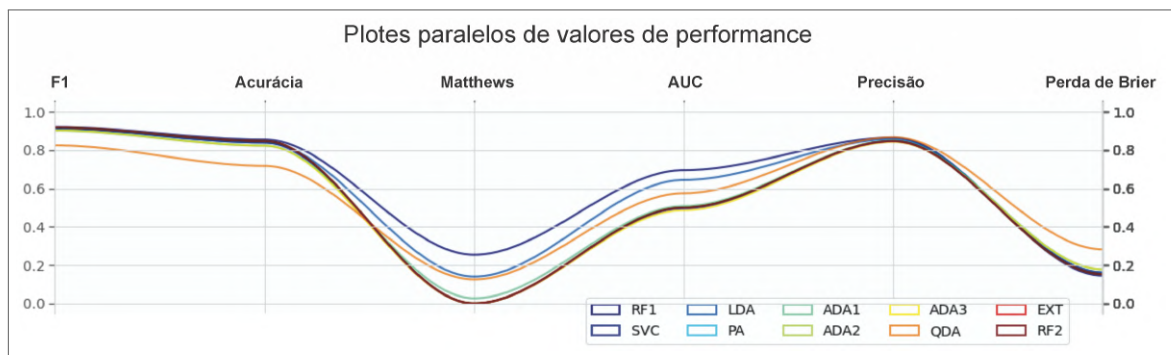
Fonte: Elaborada pelo autor (2021)

A imagem nos mostra que nem sempre a manipulação dos dados e de variáveis contribuem para o melhor desempenho do modelo de ML. Como exemplo o *Random Forest* que apresentou melhor desempenho no modelo sem nenhum tipo de pré-processamento (roxo escuro), quando comparado ao modelo com transformação polinomial dos dados e com balanceamento de classe (amarelo claro). Vale ressaltar que essa análise de forma automática através do AutoML contribui bastante para o ganho de tempo e de informação para a construção dos sistemas de predição.

Por outro lado, é importante a avaliação de modelos de ML por diferentes métricas de performance, uma vez que apenas uma métrica pode favorecer um tipo de acerto em detrimento

de outro. Isso motivou a construção da Figura 25, na qual foram adicionadas as métricas de performance F1, Matthews, área debaixo da curva (*Area Under the Curve*) (AUC), precisão e perda de brier, métricas essas disponíveis na biblioteca do matplotlib, e avaliados nos dez sistemas mais bem posicionados que foram compostos pelos classificadores: RF, SVC, LDA, análise por discriminante quadrático (*Quadratic Discriminant Analysis*) (QDA), classificador de impulso adaptativo (*Adaptive Boosting classifier*) (AdaBoost) e classificador de Árvores Extras (*Extra Trees classifier*) (ExTree).

Figura 25 – Análise dos dez sistemas com maiores acurácias fornecidos pelo Auto-sklearn por diferentes métricas de performance. Foram escolhidas as métricas comumente utilizadas para análise de modelos de ML.



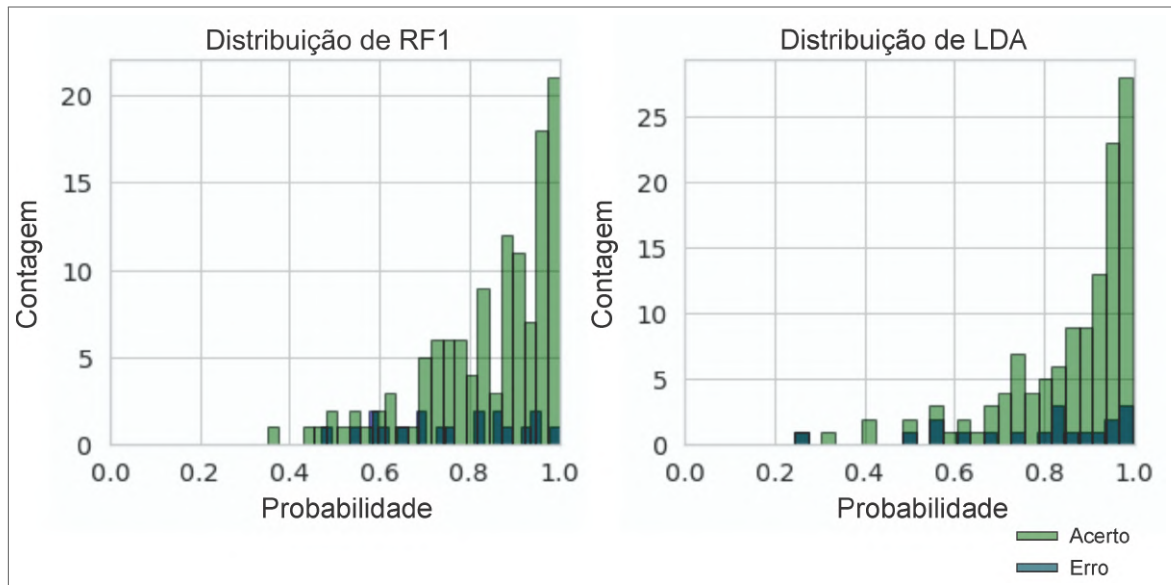
Fonte: Elaborada pelo autor (2021)

A análise em múltiplas métricas de performance permitiram a avaliação dos classificadores de modo mais específico, com métricas como a precisão e AUC que ressaltam os falsos-positivos, erro esse que é de maior relevância para os processos de garantia de qualidade. Também pode ser feita uma avaliação de performance global observando as curvas da Figura 25. De forma específicas, os modelos com melhores desempenhos foram o RF1 e SVC e o LDA, o que se repetiu para as demais métricas.

Como citado, em situações de inspeção de qualidade é importante observar o tipo de erro que os classificadores apresentam, e, principalmente, a ocorrência de falsos-positivos. A Figura 26, tem esse objetivo, busca observar a distribuição de probabilidade das predições dos dois classificadores mais bem classificados e que permitem essa avaliação, RF e LDA.

A distribuição de probabilidade de classe do modelo LDA mostra que o classificador é mais assertivo quanto a classe, por exemplo, em uma classificação correta do modelo a probabilidade da classe será elevada. Porém, o modelo de LDA também apresenta mais erros de falso-positivo com probabilidades elevadas, o que diminui a confiabilidade do modelo. Isso pode ser observado a partir da contagem de falso-positivos com probabilidade maior de 0.8, onde o modelo RF

Figura 26 – Distribuição de probabilidade de predição dos classificadores RF e LDA. A comparação entre as distribuições mostra que existe uma maior quantidade de erros de falso positivo ($p > 0.5$) para o modelo LDA.



Fonte: Elaborada pelo autor (2021)

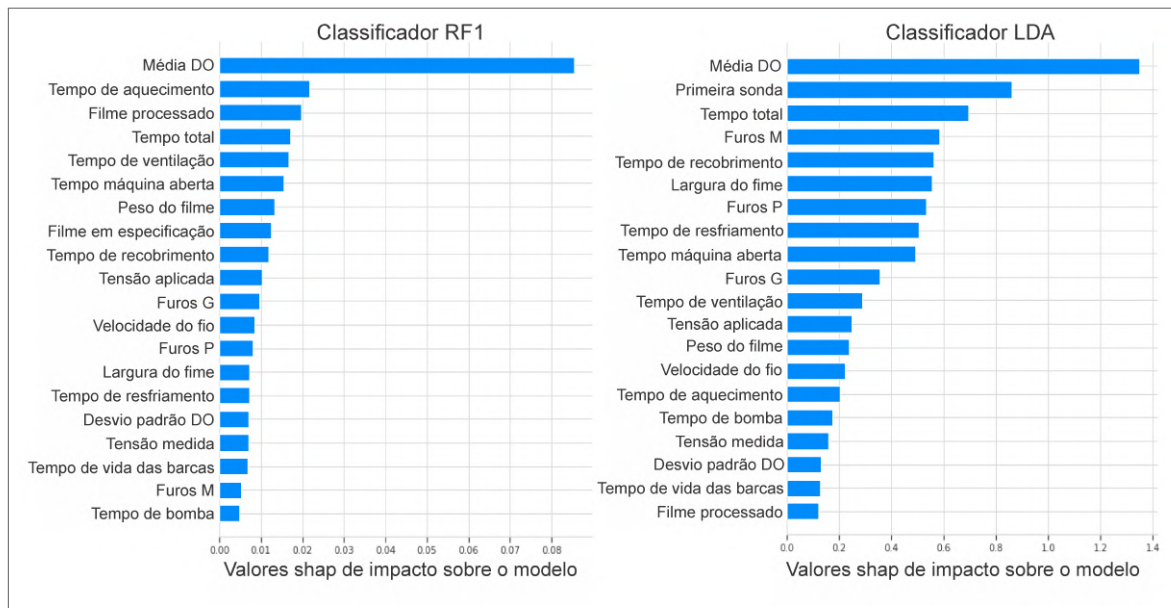
apresenta 9 erros e o modelo LDA 11 erros.

Outra forma de avaliar a predição de classe pelos modelos e através da importância das variáveis na construção dos modelos, uma vez que essas variáveis devem estar de acordo com os princípios físicos e corresponder a importância dada por especialistas. A Figura 27, mostra a relevância das variáveis para a classe do filme aprovado para os modelos RF e LDA.

É possível observar que apesar de desempenhos globais similares, devido aos diferentes modos de construção dos classificadores, os modelos RF e LDA apresentam pesos e importâncias diferentes para as variáveis de processo. Para ambos os classificadores a variável de maior importância é a média DO, algo esperado uma vez que essa variável corresponde ao valor médio de recobrimento do filme, porém, para o classificador RF essa variável apresenta maior relevância do que para o classificador LDA. Essa diferença na distribuição de pesos revela maior sensibilidade do classificador RF a variável Média DO. Apesar das demais variáveis apresentarem pesos próximos, a hierarquia de importância das variáveis obtida pelo classificador RF demonstra ser mais próxima da avaliação de especialistas, uma vez que ela fornece maiores pesos a propriedades mais sensíveis, como média DO e tempo de aquecimento.

Uma vez que é conhecido o impacto das variáveis sobre o classificador de melhor desempenho, é importante analisar quais valores de variáveis favorecem a uma dada classe, e se isso condiz com a realidade do processo. Pensando nisso, a Figura 28 mostra a correlação entre a

Figura 27 – Importância das variáveis para os classificadores RF1 e LDA, de acordo com os valores SHAP médios de impacto. O que permite a observância entre a importância real das variáveis do processo industrial e como elas afetam no processo de predição de qualidade.



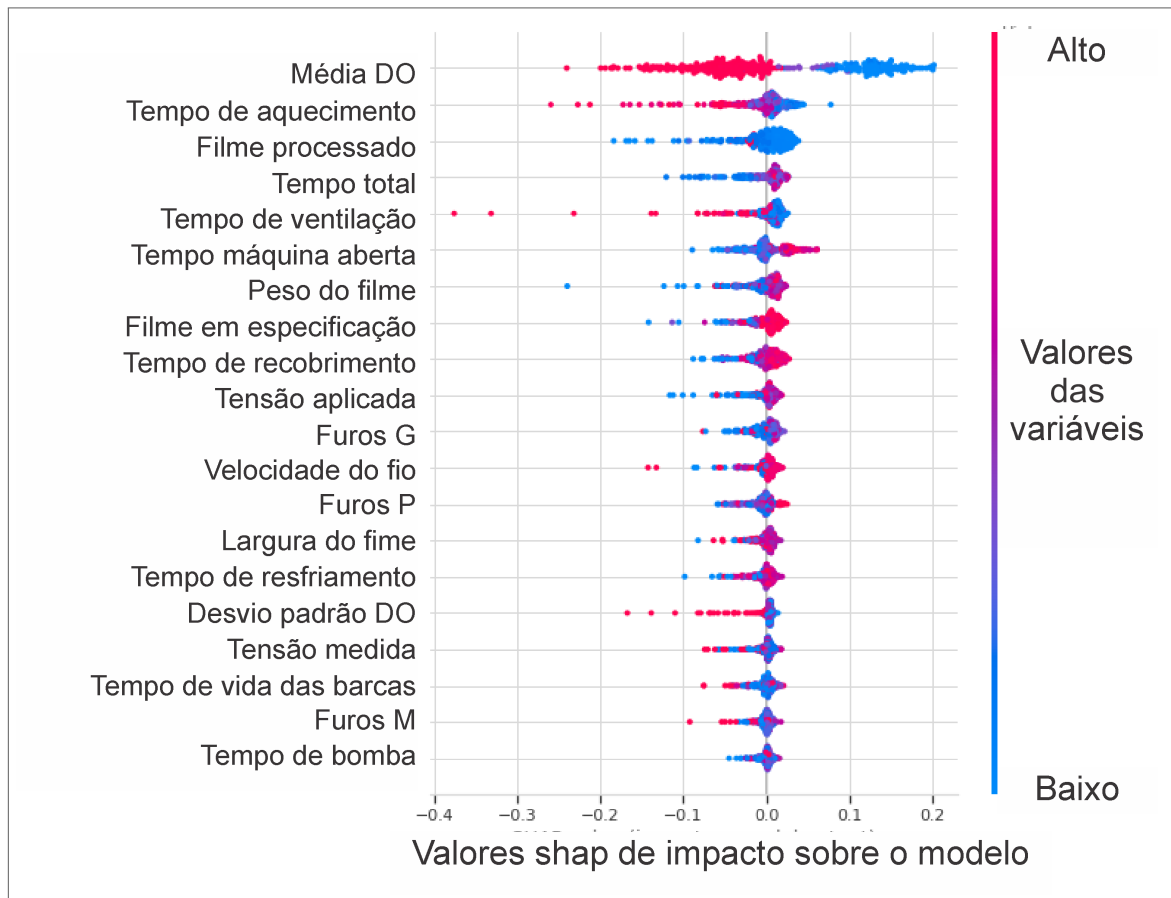
Fonte: Elaborada pelo autor (2021)

distribuição de valores das variáveis do problema com o impacto gerado sobre a classificação final do produto.

A dispersão dos valores da variável Média DO, explica bem o porque dela ser a variável de maior peso sobre os modelos com melhores desempenhos, é possível observar que de valores altos dessa variável impactam negativamente na aprovação do produto, já os valores mais baixos contribuem para a classe de filmes aprovados. Seguindo o princípio de recobrimento em superfícies de baixa energia mostrado na Figura 3, leva a entender que o fator de Média DO ser negativo para o sistema se deve ao preenchimento não uniforme da superfície do filme, o que leva a regiões com acúmulo de metais, aumentando o valor médio de OD, porém com um recobrimento de forma irregular, levando a baixa qualidade de recobrimento e possível reprovação do produto formado.

A segunda variável mais importante para o modelo RF, o tempo de aquecimento das barcas evaporadoras, demonstra um comportamento similar a média de OD, no qual valores mais altos representam um impacto negativo na aprovação do produto. Diferentemente da média DO, os valores de tempo de aquecimento devem variar de acordo com o tempo de vida das barcas evaporadoras, uma vez que, barcas mais novas (menor tempo de vida) apresentam melhor dispersão dos componentes da barca e eliminação de impurezas, através de um maior tempo de aquecimento, diferentemente de barcas mais desgastadas, na qual o aquecimento prolongado

Figura 28 – Dispersão dos valores das variáveis do processo que compõem os objetos, e os seus respectivos impactos sobre o modelo RF para predição de filmes aprovados. Demonstrando a correlação existente entre a discriminação dos valores das variáveis e a sua significância para a construção do modelo.



Fonte: Elaborada pelo autor (2021)

pode levar a fadiga do material e a dispersão desuniforme dos componentes da barca, levando a perda de propriedades como a resistividade. Assim, o uso de valores mais extremos de tempo de aquecimento podem impactar significativamente na qualidade do produto obtido.

Para as demais variáveis foi possível observar que, em geral, elas representam baixo impacto sobre o modelo, a menos de alguns poucos exemplos onde valores mais extremos para essas variáveis representaram de forma significativa um impacto para reprovação do produto. O que leva a considerar que apesar dessas variáveis não apresentaram em média grande impactos para o classificador, o descontrole de seus valores podem representar um prejuízo a qualidade final do filme, favorecendo a reprovação desses produtos. Cabendo assim também, a importância e o monitoramento dessas variáveis durante o processo de manufatura.

Com isso, é possível observar que além da obtenção de um modelo de RF com boa performance em diversas métricas de desempenho, o mapeamento e a distribuição de pesos das

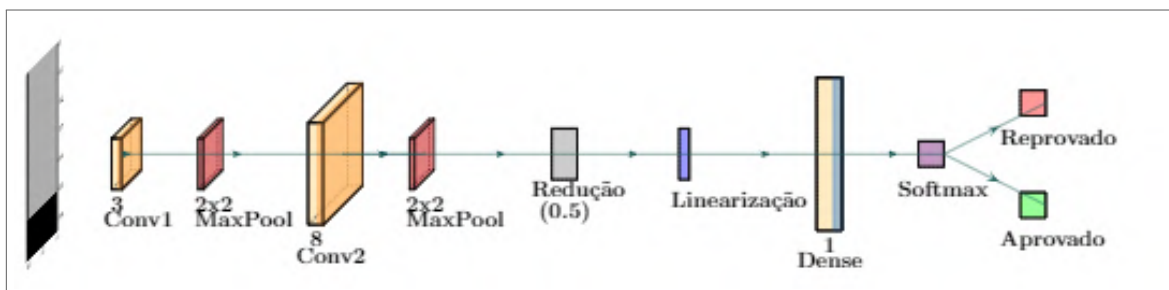
variáveis do problema, são condizentes com a realidade do processo e a dinâmica físico-química que regem o comportamento dessas variáveis.

4.2 ANÁLISE DOS FILMES POR VISUALIZAÇÃO DO PERFIL DE DO

Como visto na seção 2.4, modelos de processamento de imagens contribuem com a extração de padrões que podem auxiliar para o entendimento de uma imagem. Partindo disso, foi pensado na utilização dos dados da Tabela 2 para construção de imagens de 21x65, 21 sensores com 65 medidas de OD, utilizando os perfis de OD dos produtos obtidos.

Essas imagens foram utilizadas tanto como informação para predição de classe do produto, como aprovado ou reprovado, quanto para aplicação de modelo não-supervisionado de segmentação semântica. A Figura 29, mostra o modelo de rede neural proposto para classificação dos filmes a partir dos dados de OD.

Figura 29 – Modelo de rede neural sequencial do Keras adaptada para classificação dos filmes metalizados utilizando informações dos perfis de OD. O modelo é composto por duas camadas convolucionais alternadas de duas camadas de conjugação 2x2. Em seguida é aplicada uma regularização de 0.5 para evitar sobreajuste do modelo, e em seguida a linearização e determinação da classe dos objetos pela função de ativação softmax.



Fonte: Elaborada pelo autor (2021)

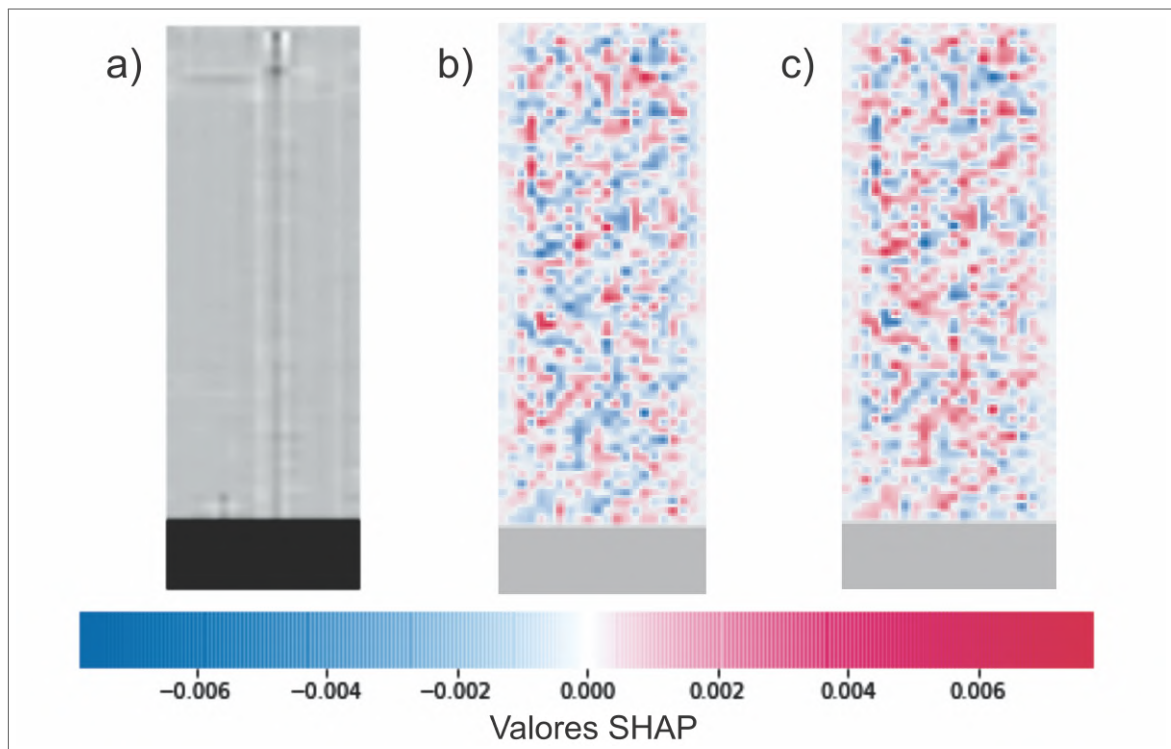
A rede neural proposta foi construída utilizando a biblioteca Keras¹ com determinação de sua estrutura através de experimentação prévia, e ao final representada de forma sequencial por uma camada convolucional com 3 neurônios e função de ativação ReLU, seguida por um *max-pooling* de 2x2 e outra camada convolucional de 8 neurônios com função de ativação ReLU também seguida por um *max-pooling* de 2x2, para evitar o sobre-ajuste do modelo, foi feita uma exclusão de 50% dos pesos. Após isso, os dados foram linearizados através do método *flattening* junto a uma camada de rede densa de uma dimensão, e posterior aplicação da função *softmax* para a classificação final.

¹ <https://keras.io/>

Esse modelo final foi obtido através de experimentação observando valores de performance. O modelo proposto apresentou acurácia de 86,67% e perda de 0.46. A performance obtida se justifica pelo número limitado de objetos presentes no banco de dados, aliado a uma elevada quantidade de informação para mapeamento.

Para analisar o processo de predição do modelo, foram observados os valores de SHAP para as classes aprovados e reprovados, como forma de identificar se o modelo está mapeando de forma correta as variáveis para as classes (Figura 30). Os valores de SHAP representam o impacto das variáveis, indicando se os seus valores favorecem ou não uma determinada classe.

Figura 30 – a) Perfil de OD do filme metalizado. Valores SHAP de impacto sobre a predição das classes b) Aprovado e c) Reprovado. Os valores positivos de SHAP, marcados em vermelho, favorecem a classe observada, já os valores em azul prejudicam a classificação observada. O que permite a visualização da distribuição desses valores de SHAP de acordo com a qualidade do recobrimento.



Fonte: Elaborada pelo autor (2021)

Através da Figura 30a, foi possível observar que o perfil do filme apresenta variações nos valores de OD em diferentes posições de recobrimento, principalmente, na região central do filme, representando prováveis problemas nos sistemas de recobrimento desses locais. Os valores de SHAP nessa região permitem observar um desfavorecimento da classe aprovado (Figura 30b) e favorecendo a classe reprovado (Figura 30c). Como a alta variabilidade dos valores de OD prejudica a qualidade do recobrimento era esperado que de fato os valores de SHAP favorecem a classe reprovado, demonstrando o correto mapeamento das variáveis pela

modelo de rede neural construído.

A análise dos valores de SHAP também contribuem para a observação de regiões com possíveis falhas de recobrimento, os quais em regiões com padrões verticais de falha podem ser entendidas como problemas locais no sistema de recobrimento. Isso tende a contribuir para uma análise imediata do recobrimento e a tomada de decisão para ajustes de configurações e de condições de variáveis e equipamentos do processo.

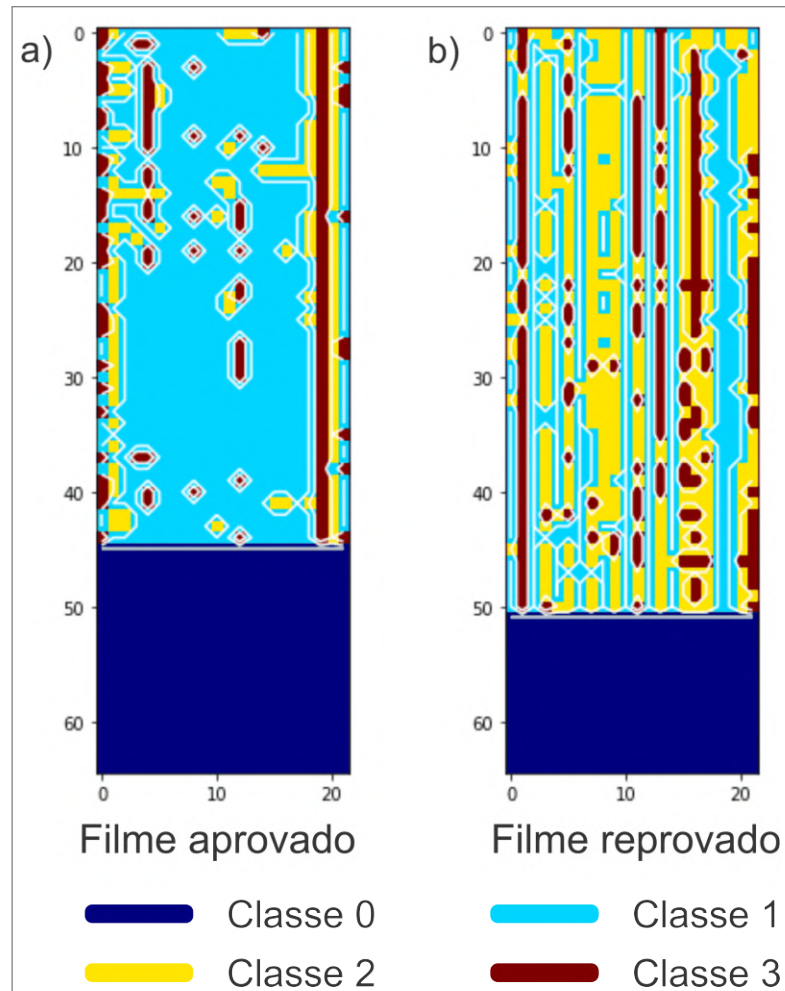
Com o mesmo entendimento de análise do perfil de OD, a Figura 31 mostra a segmentação obtida em dois diferentes objetos, o primeiro é o perfil de OD segmentado de um filme aprovado (Figura 31a), e o segundo de um filme reprovado (Figura 31b), nos quais cada cor indica uma classe a qual a região foi endereçada pelo modelo de segmentação de imagens em classes.

A segmentação semântica dos objetos permite identificar padrões a serem correlacionados a possíveis falhas e defeitos presentes no filme, sendo a classe 0 representante da região da imagem sem a presença de filme. A Figura 31a é formada majoritariamente pela classe 1, azul claro, e apresenta pequenas regiões endereçadas com a classe 2, amarelo, e na posição vertical 19 uma área marcada por vermelho, classe 3. Já a Figura 31b, apresenta majoritariamente a classe 2, com a classe 1 em menor quantidade quando comparada com a área apresentada no filme aprovado, mostrando também diversas posições verticais do recobrimento marcadas em vermelho.

Falhas comuns presentes nos recobrimentos metálicos de filmes são ranhuras, arranhões, pequenos furos e falhas dos sistemas locais que geram uma cobertura de baixa qualidade por todo o comprimento do filme na sua respectiva posição. Pelo formato das regiões obtidas, é possível correlacionar essas falhas a classe 3 identificadas em maiores quantidades para o filme reprovado. A classe 3 apresenta formatos de pontos, traços e colunas, que podem ser interpretados como furos, arranhões e falhas do sistemas de recobrimento, respectivamente. Essas falhas do sistema de recobrimento, seja por velocidade ou alinhamento do fio metálico, ou mesmo pela temperatura da barca evaporadora, nas posições 2, 11, 13, 16 e 22 da Figura 31b levam a alta variabilidade de OD e baixa qualidade de filme. Já os filmes aprovados apresentam em sua maioria região azul, classe 1, mostrando maior uniformidade no recobrimento, a menos de regiões espaços em vermelho, que podem ser identificadas como furos ou falhas de recobrimento causadas pela presença de sujeiras na superfície do filme não recoberto, mas que não são suficientes para a reprovação do produto.

Para ajudar na observação do mapeamento de classes e sua correlação com a qualidade do recobrimento do filme, foi feita a Figura 32. Onde cada gráfico demonstra a proporção de

Figura 31 – Segmentação semântica em quatro classes dos perfis de OD dos filmes para identificação de falhas. As diferentes cores representam as classes obtidas por agrupamento, e a área em azul marinho representa área onde não há filme. De acordo com o formato e a distribuição das classes, é possível inferir que a classe 1, azul claro, é uma provável região de bom recobrimento, a classe 2, amarelo, é uma provável área de recobrimento intermediário, e por fim a classe 3, em vermelho, como representante de regiões de recobrimento pobre. a) Perfil segmentado de um objeto aprovado. b) Perfil segmentado de um objeto reprovado.

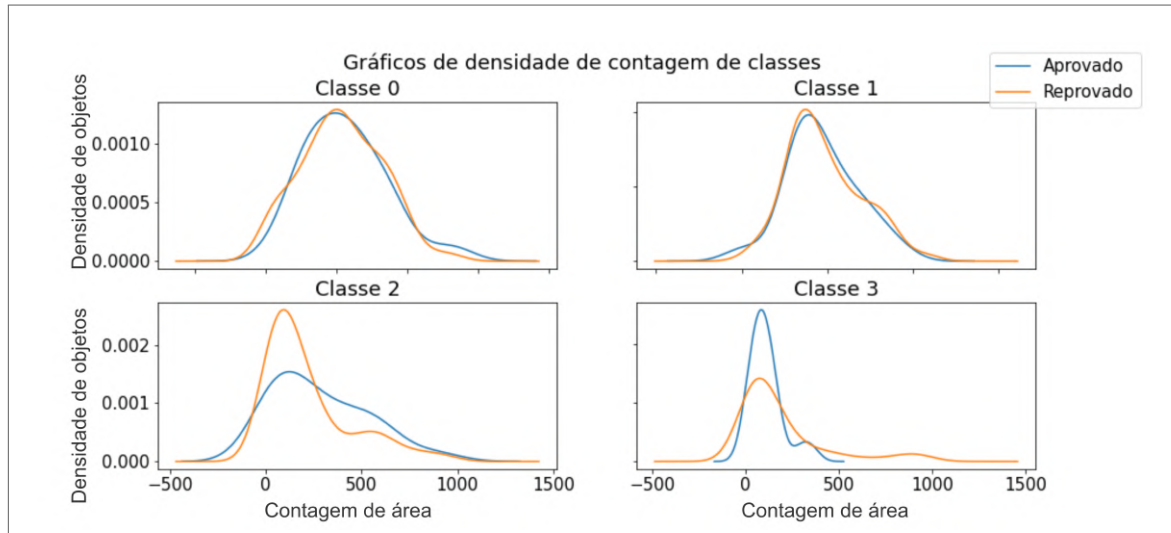


Fonte: Elaborada pelo autor (2021)

objetos de acordo com a contagem de área da respectiva classe.

Essa figura mostra em particular que os exemplos reprovados, apresentaram alta contagem de área para a classe 3, corroborando com a ideia de que essa classe representa falhas no recobrimento e que elas contribuem para a reprovação do produto. Outra observação que pode ser feita é que apesar da classe 2 poder ser entendida como uma classe intermediária entre um bom recobrimento e a ocorrência de falhas, ela não é uma classe determinante de qualidade uma vez que não apresenta um padrão de contagens altas para os filmes reprovados. Porém, essa classe pode ser levada em consideração como uma classe indicadora de regiões de potenciais falhas de recobrimento que merecem atenção, e se necessário, ajustes do sistema

Figura 32 – Demonstração dos valores de contagem de acordo com as classes utilizadas na segmentação semântica dos perfis de OD dos filmes. As barras pretas verticais demonstram os valores de contagem para os diferentes exemplos cuja classificação é descrita de acordo com a cor da barra, verde para filmes aprovados e amarelo para filmes reprovados.



Fonte: Elaborada pelo autor (2021)

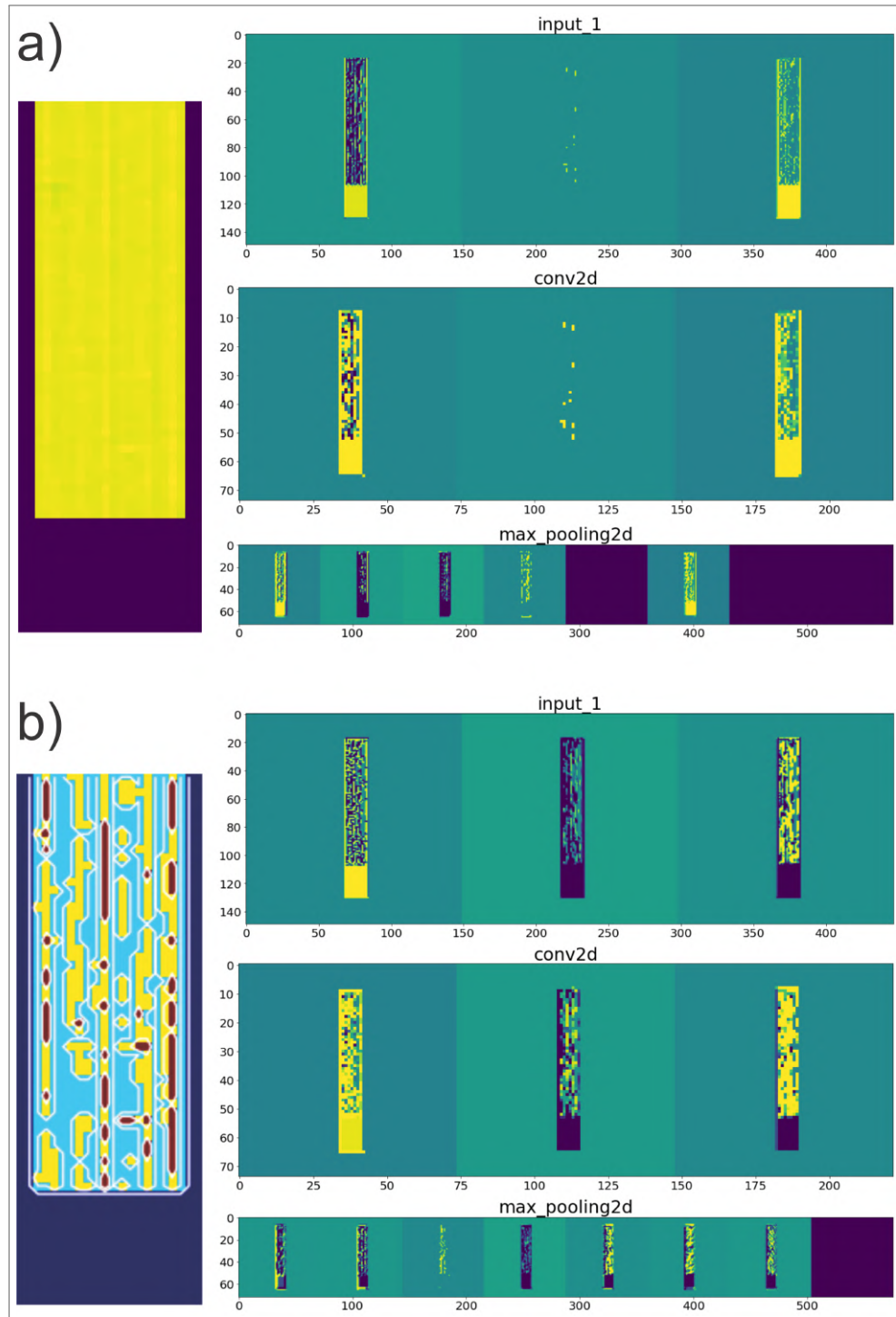
local de recobrimento.

Uma vez que a segmentação semântica do perfil dos objetos além de dividir a imagem em diferentes classes, constrói um contorno para essas diferentes classes e regiões, ela também pode ser explorada quanto a contribuição para melhores desempenhos na predição de classe dos objetos, já que modelos de redes neurais apresentam bom desempenho no mapeamento de contornos e segmentos.

Utilizando o mesmo modelo de rede neural presente na Figura 29, os perfis de OD segmentados, foram utilizados para a predição de classe dos objetos do conjunto de validação. A Figura 33, demonstra a comparação entre as saídas das 3 primeiras camadas do modelo de rede neural, tanto para o modelo treinado com os perfis de OD originais (Figura 33a), quanto para o modelo treinado com os perfis de OD após segmentação semântica (Figura 33b).

A comparação das saídas obtidas entre o treinamento do classificador com perfis de OD segmentados e não segmentados demonstra um importante ganho do mapeamento de informações por parte do modelo de ML aplicado. O primeiro ganho observado é na distribuição de informação, a qual o mapeamento feito diretamente do perfil de OD original apresenta um dos neurônios da primeira camada com pesos praticamente nulos, ou com baixa quantidade de informação, o que é propagado para as camadas subsequentes fazendo com que o modelo treinado com os objetos originais apresente na 3ª camada, três neurônios vazios de informação, enquanto que o modelo treinado com os objetos segmentados apresenta apenas um neurônio

Figura 33 – Análise de sensibilidade e espacialidade de informações a partir das saídas das três primeiras camadas do modelo de rede neural treinados com imagens dos perfis de OD originais (a esquerda) e com os perfis após segmentação semântica (a direita). a) Mapeamento de variáveis utilizando objetos com perfis de OD originais. b) Mapeamento de variáveis utilizando objetos com perfis de OD após segmentação semântica.



Fonte: Elaborada pelo autor (2021)

vazio de informação para a mesma camada.

O segundo ganho observado no modelo treinado sobre os objetos segmentados é quanto

a sensibilidade do modelo. A Figura 33b apresenta maiores níveis de detalhamento nos pesos atribuídos as variáveis, o que contribui também na detecção e influência de falhas para a classificação final do algoritmo. O modelo treinado com os objetos segmentados apresentou desempenho semelhante ao modelo treinado com os perfis de OD originais, 86,67% e perda de 0.6, o que pode ser entendido pelo baixo número de objetos para o processo de treinamento e validação. Por outro lado, o ganho de espacialidade e de sensibilidade do modelo revelam a importância do treinamento do modelo com os objetos segmentados, além da contribuição para identificação e visualização de falhas. Para trabalhos futuros, a segmentação semântica pode ser comparada com a análise real do filme por especialistas e equipamentos, a fim de comprovar a associação das classes com a qualidade do recobrimento.

A aplicação dos modelos de processamento de imagem aliados a análise de especialistas e operadores, podem contribuir para avaliação e seleção de objetos para posterior análises laboratoriais, reduzindo a aleatoriedade da seleção e os custos associados a perda de materiais e de tempo com as análises laboratoriais. Além de auxiliar nos processos de monitoramento e manutenção de equipamentos e variáveis envolvidas no processo de deposição metálica.

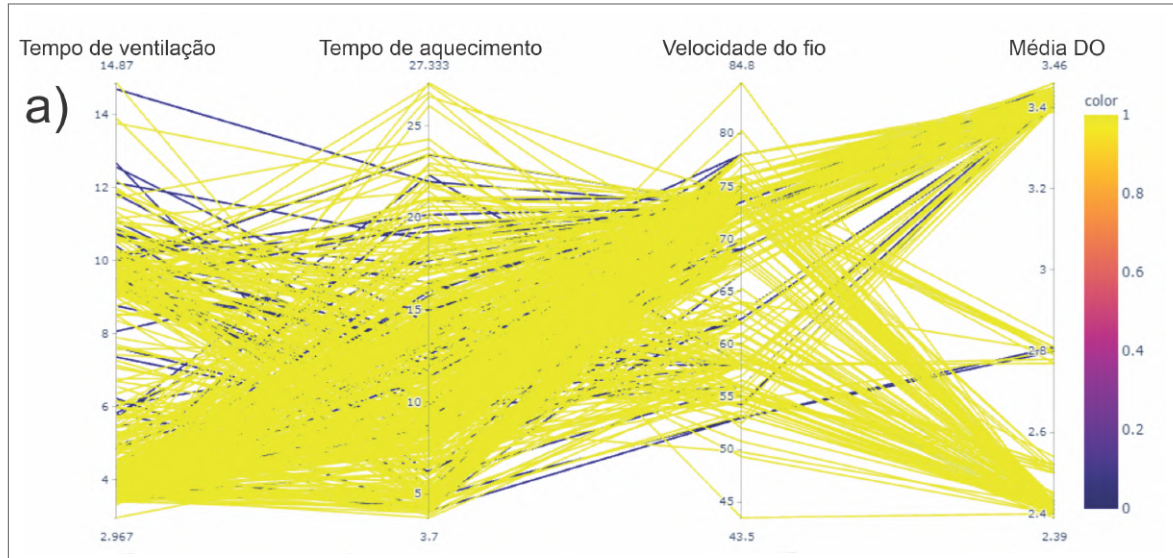
4.3 INTERPRETAÇÃO DE VARIÁVEIS DE PROCESSO E DO PRODUTO

De forma complementar a utilização dos modelos de ML e a sua análise na predição da aprovação ou não dos filmes metalizados, faz-se importante também a análise das variáveis com maiores importâncias para a qualidade do processo e os seus impactos, já que são essas que determinam a qualidade final do produto e também a predição obtida pelo modelo de ML. Nesse intuito, foram feitas aplicações para visualização de objetos com a finalidade de resumir informações operacionais e dos modelos de ML aplicados.

O primeiro modelo de visualização é um plote de eixos paralelos contendo as quatro principais variáveis operacionais para o processo de predição de classe: tempo de ventilação, tempo de aquecimento, velocidade do fio e média DO. A Figura 34 mostra o histórico de condições operacionais para filmes aprovados, em linhas amarelas, e reprovados, em linhas azuis.

Essa visualização permite a observação das correlações existentes entre as variáveis vizinhas, como exemplo o tempo de aquecimento e a velocidade do fio que apresentam correlações positivas. Isso permite a análise de operadores para o melhor controle operacional e ajuste das variáveis em tempo de produção. Já a Figura 35 demonstra os mesmos dados históricos limitados a certas condições operacionais, o que nos permite observar a proporção de produtos

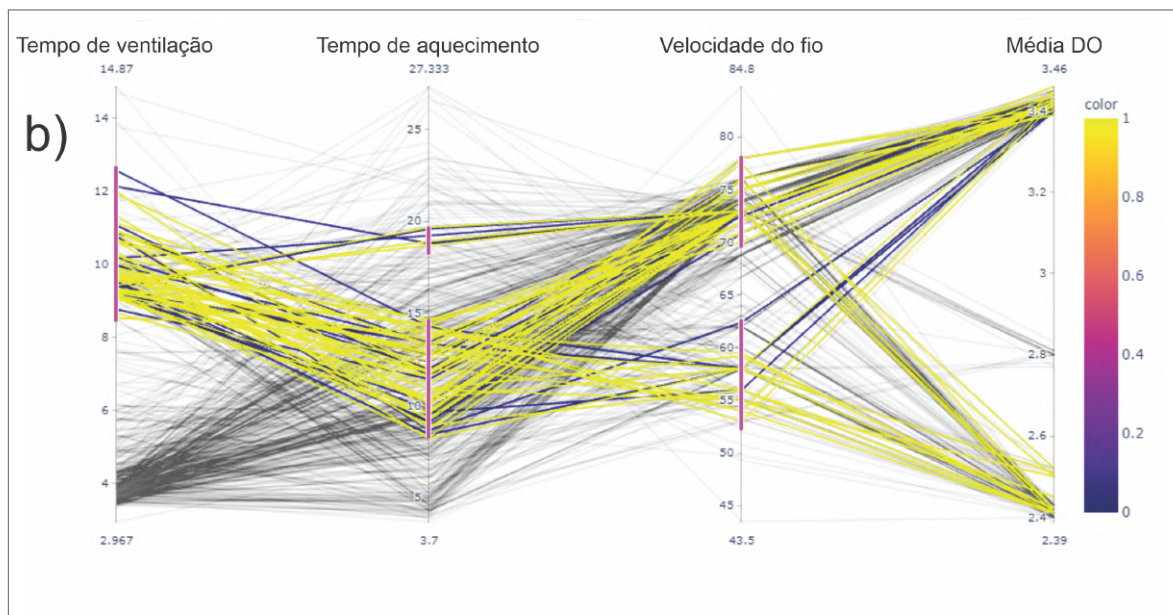
Figura 34 – Utilização de plotes paralelos com variáveis de processo mais relevantes para a predição pelo modelo RF. O que permite a observação do histórico de correlação das variáveis para objetos aprovados, linhas amarelas, e para objetos reprovados, linhas azuis. Mostrando a existência de correlações diretas e correlações inversas entre variáveis vizinhas, o que pode auxiliar operadores na escolha de parâmetros operacionais que favoreçam maiores taxas de produtos aprovados.



Fonte: Elaborada pelo autor (2021)

reprovados ao serem adotadas certas condições operacionais.

Figura 35 – Utilização de plotes paralelos com valores de variáveis de processo mais relevantes restritos a certas condições operacionais. Demonstrando que há influência da combinação de valores das variáveis de processo sobre a taxa de produtos aprovados. E que a certas condições operacionais devem ser evitadas, já que apresentam uma elevada taxa de produtos reprovados.



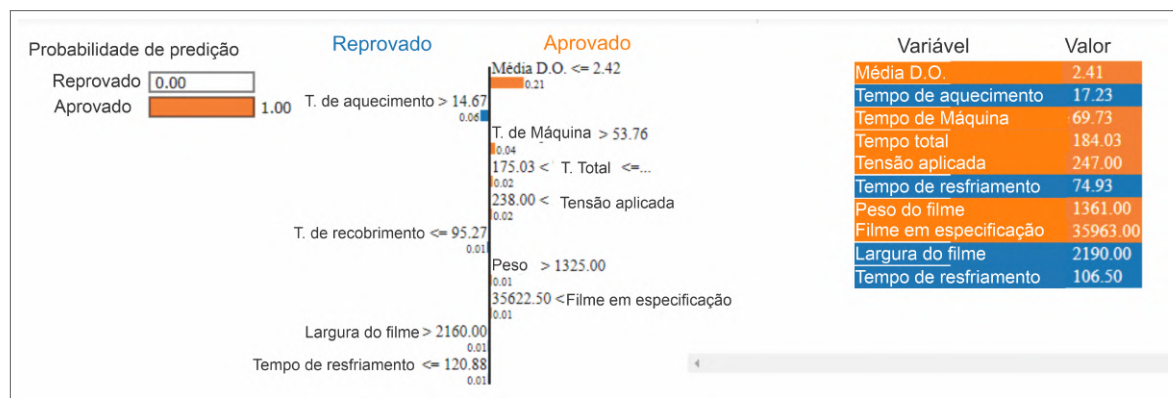
Fonte: Elaborada pelo autor (2021)

É possível observar através da comparação com a Figura 34 que determinadas condições

operacionais aumentam a proporção de produtos reprovados, demonstrando que além do controle de variáveis, é importante a análise da combinação presente entre elas. Por exemplo, conforme mostrado, para a produção de filmes com tempo de ventilação de 9 a 12 min e tempo de aquecimento de 10 a 15 min, a utilização da velocidade do fio na faixa de 70 m/-min tende a produzir uma maior taxa de produtos reprovados, mostrando para operadores e analistas que essa combinação de condições operacionais deva ser evitada.

A segunda forma de visualização de objetos retrata um resumo dos valores operacionais e de seus impactos sobre a probabilidade final de aprovação e reprovação de um determinado objeto. A Figura 36 mostra um objeto aprovado com as informações operacionais utilizadas na sua produção e os impactos gerados por eles.

Figura 36 – Resumo de predição de classe do modelo de ML para um exemplo aprovado. Essa visualização demonstra os impactos das variáveis sobre a classificação do produto, e a descrição das condições operacionais durante a sua manufatura. Tendo como objetivo, permitir operadores e analistas a observação e interpretação dos produtos e da sua classificação em tempo de produção.



Fonte: Elaborada pelo autor (2021)

A visualização acima mostra como a variável Média OD impacta positivamente na aprovação do objeto, além do impacto negativo sobre a aprovação pelas variáveis de tempo de aquecimento e de recobrimento, fornecendo a observação das condições operacionais que precisam ou não de ajustes, guiando assim operadores e analistas nos processos de monitoramento e manutenção.

Os softwares desenvolvidos tanto para a predição de variáveis quanto para a visualização dos filmes e das condições operacionais servem como base para uma implementação de monitoramento em tempo de produção. Como esses dados de produção, atualmente, já são coletados e armazenados em disco rígido, poderiam ser armazenados em um disco compartilhado via rede e exportado para um programa em nuvem, onde pode ser utilizado por alguma plataforma que compile os códigos e processe os dados em tempo de produção, como google

Colab[®].

5 CONCLUSÃO

Nesse trabalho foi possível explorar diferentes tipos de classificadores de ML, associados a diferentes condições de pré-processamento de dados e de variáveis para o processo de predição qualidade dos produtos obtidos no processo industrial de metalização a vácuo, tendo o modelo *Random Forest* o melhor desempenho com 85,4% de acurácia. Por outro lado, através da segmentação semântica dos perfis de OD dos filmes foi possível a identificação de falhas e o monitoramento da qualidade dos filmes produzidos. A análise dos perfis de OD segmentados propiciaram também melhorias na sensibilidade e espacialidade do modelo de rede neural aplicado sobre os valores de OD para predição de classe do produto, apresentando acurácia de 86,67% e perda de 0,6.

A exploração de modelos de ML forneceu predições de variáveis que podem vir a ser implementadas de forma complementar as análises laboratoriais que são atualmente feitas nos filmes metalizados através pequenos recortes do filme, reduzindo a possibilidade erros de inspeção devido à variabilidade espacial presente no recobrimento metálico. Essa combinação de inspeções de qualidade tende a promover a redução de perda de material por recortes e possíveis erros por generalização. De forma complementar, foram aplicadas visualizações que forneceram o entendimento dos produtos obtidos, revelando correlações importantes entre as variáveis e as condições operacionais, e como impactam sobre a aprovação dos filmes metalizados. O que antes não era levado em consideração para o processo produtivo de filmes metalizados, os quais tinham seus parâmetros ajustados de forma reativa e sem levar em conta os dados históricos.

Apesar da limitação quantitativa de informações, os sistemas explorados se mostraram capazes de oferecer suporte a analistas e operadores para o monitoramento e a manutenção inteligente do processo de metalização à vácuo, cabendo o aprimoramento dos modelos de maiores desempenhos para serem implementados em produção. A exploração dos modelos de ML e as aplicações desenvolvidas nesse trabalho abrem portas para implementações futuras que utilizem a inteligência artificial em processos industriais. Tornando o processo de metalização a vácuo ainda mais preciso, com integração de dados e de informações característicos dos processos de manufatura avançada, além de demonstrar o potencial de melhoria de performance que as indústrias 4.0 podem oferecer.

REFERÊNCIAS

- AGARAP, A. F. Deep learning using rectified linear units (relu). *CoRR*, abs/1803.08375, 2018. Disponível em: <<http://arxiv.org/abs/1803.08375>>.
- AKBARI, M.; LIANG, J.; HAN, J.; TU, C. Learned multi-resolution variable-rate image compression with octave-based residual blocks. *arXiv preprint arXiv:2012.15463*, 2020.
- AMERSHI, S.; CHICKERING, M.; DRUCKER, S. M.; LEE, B.; SIMARD, P. Y.; SUH, J. ModelTracker: Redesigning Performance Analysis Tools for Machine Learning. In: *Proceedings of the 33rd ACM SIGCHI Conference on Human Factors in Computing Systems*. [S.l.]: ACM, 2015. p. 337–346. ISBN 978-1-4503-3145-6.
- AMR, T. *Hands-On Machine Learning with scikit-learn and Scientific Python Toolkits: A practical guide to implementing supervised and unsupervised machine learning algorithms in Python*. Packt Publishing, 2020. ISBN 9781838823580. Disponível em: <<https://books.google.nl/books?id=GlbzDwAAQBAJ>>.
- ARCHIBALD, P.; PARENT, E. Source evaporant systems for thermal evaporation. *Solid State Technology*, IEEE, p. 32–40, 1976.
- BARREIRO, E.; MUNTEANU, C. R.; CRUZ-MONTEAGUDO, M.; PAZOS, A.; GONZÁLEZ-DÍAZ, H. Net-net auto machine learning (automl) prediction of complex ecosystems. *Scientific Reports - Nature*, Springer, v. 8, 2018.
- BELKACEM, S. *Machine learning approaches to rank news feed updates on social media*. Tese (Theses) — Université des Sciences et de la Technologie Houari Boumediene Alger, abr. 2021. Disponível em: <<https://tel.archives-ouvertes.fr/tel-03201363>>.
- BERGSTRA, J.; BREULEUX, O.; BASTIEN, F.; LAMBLIN, P.; PASCANU, R.; DESJARDINS, G.; TURIAN, J.; WARDE-FARLEY, D.; BENGIO, Y. Theano: A cpu and gpu math compiler in python. In: *Proceedings of the 9th Python in Science Conference*. [S.l.: s.n.], 2010. p. 3–10.
- BERGSTRA, J.; YAMINS, D.; COX, D. D. et al. Hyperopt: A python library for optimizing the hyperparameters of machine learning algorithms. In: CITESEER. *Proceedings of the 12th Python in science conference*. [S.l.], 2013. v. 13, p. 20.
- BIJI, K. B.; RAVISHANKAR, C. N.; MOHAN, C. O.; GOPAL, T. K. S. Smart packaging systems for food applications: a review. *Journal of Food Science and Technology*, Springer, v. 52, n. 10, 2015.
- BIKMUKHAMETOV, T.; JÄSCHKE, J. Combining machine learning and process engineering physics towards enhanced accuracy and explainability of data-driven models. *Computers and Chemical Engineering*, Elsevier, v. 238, 2020.
- BISHOP, C. A. *Vacuum deposition onto webs, films and foils 1st Edition*. [S.l.]: C.A. Bishop Consulting Ltd United Kingdom and EMMOUNT Technologies Canandaigua NY United States, 2007. ISBN 0-8155-1535-9.
- BISHOP, C. A. *Vacuum deposition onto webs, films and foils 2nd Edition*. [S.l.]: ohn Wiley & Sons, Inc. Hoboken, New Jersey, and Scrivener Publishing LLC, Salem, Massachusetts., 2011.

BISHOP, C. A. *Roll-to-Roll Vacuum Deposition of Barrier Coatings 2nd Edition*. [S.l.]: John Wiley & Sons, Inc. Hoboken, New Jersey, and Scrivener Publishing LLC, Salem, Massachusetts., 2015.

BISHOP, C. A.; III, E. M. M. Vacuum metallizing for flexible packaging. *MULTILAYER FLEXIBLE PACKAGING*, C.A. Bishop Consulting Ltd United Kingdom and EMMOUNT Technologies Canandaigua NY United States, p. 235–255, 2016.

BISHOP, C. A.; Mount III, E. M. *Vacuum Metallizing for Flexible Packaging*. [S.l.]: Ohn Wiley & Sons , Inc. Hoboken, New Jersey, and Scrivener Publishing LLC, Salem, Massachusetts., 2012. v. 5.

BISHOP, C. M. Bayesian pca. *Advances in neural information processing systems*, MIT; 1998, p. 382–388, 1999.

BOROVEC, J.; KYBIC, J.; NAVA, R. Detection and localization of drosophila egg chambers in microscopy images. In: WANG, Q.; SHI, Y.; SUK, H.-I.; SUZUKI, K. (Ed.). *Machine Learning in Medical Imaging*. Cham: Springer International Publishing, 2017. p. 19–26.

BOROVEC, J.; ŠVIHLÍK, J.; KYBIC, J.; M.D., D. H. Supervised and unsupervised segmentation using superpixels, model estimation, and graph cut. *Journal of Electronic Imaging*, SPIE, v. 26, n. 6, p. 1 – 17, 2017. Disponível em: <<https://doi.org/10.1117/1.JEI.26.6.061610>>.

BOYKOV, Y.; VEKSLER, O.; ZABIH, R. Fast approximate energy minimization via graph cuts. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, v. 23, n. 11, p. 1222–1239, 2001.

CAI, H.; CHEN, T.; ZHANG, W.; YU, Y.; WANG, J. Efficient architecture search by network transformation. Association for the Advancement of Artificial Intelligence, 2018.

CHEN, Y.; MENG, G.; ZHANG, Q.; XIANG, S.; HUANG, C.; MU, L.; WANG, X. Renas: Reinforced evolutionary neural architecture search. 2019.

CHEN, Z.; NIU, J.; LIU, X.; TANG, S. Selectscale: Mining more patterns from images via selective and soft dropout. *arXiv preprint arXiv:2012.15766*, 2020.

CHOUDHARY, P.; KRAMER, A.; TEAM, c. datascience.com. *Skater: Model Interpretation Library*. 2018. Disponível em: <<https://doi.org/10.5281/zenodo.1198885>>.

COLEGROVE, L. Artificial intelligence in the chemical industry—why my industry puzzles over our vendors' struggles. *Journal of Advanced Manufacturing and Processing*, v. 2, n. 3, p. e10052, 2020. Disponível em: <<https://aiche.onlinelibrary.wiley.com/doi/abs/10.1002/amp2.10052>>.

COLLOBERT, R.; KAVUKCUOGLU, K. Torch7: A matlab-like environment for machine learning. In: *In BigLearn, NIPS Workshop*. [S.l.: s.n.], 2011.

CORTES, C.; VAPNIK, V. Support-vector networks. *Machine learning*, Springer, v. 20, n. 3, p. 273–297, 1995.

DANGETI, P. *Statistics for Machine Learning*. [S.l.]: Packt, 2017.

- DAOUTIDIS, P.; LEE, J. H.; HARJUNKOSKI, I.; SKOGESTAD, S.; BALDEA, M.; GEORGAKIS, C. Integrating operations and control: A perspective and roadmap for future research. *Computers and Chemical Engineering*, p. 179–184, 2018.
- DAVISA, J.; EDGAR, T.; PORTER, J.; BERNADEND, J.; SARLIE, M. Smart manufacturing, manufacturing intelligence and demand-dynamic performance. *Computers and Chemical Engineering*, Elsevier Ltd, v. 47, p. 145–156, 2012.
- EIGEN, D.; FERGUS, R. *Predicting Depth, Surface Normals and Semantic Labels with a Common Multi-Scale Convolutional Architecture*. 2015.
- F., M.; M., W.; E., H. *ML-plan: Automated machine learning via hierarchical planning*. *Machine Learning*, Springer, v. 107, 2018.
- FEURER, M.; EGGENSPERGER, K.; FALKNER, S.; LINDAUER, M.; HUTTER, F. *Auto-Sklearn 2.0: The Next Generation*. 2020.
- FEURER, M.; KLEIN, A.; EGGENSPERGER, K.; SPRINGENBERG, J.; BLUM, M.; HUTTER, F. Efficient and robust automated machine learning. URL <http://papers.nips.cc/paper/5872-efficient-and-robust-automated-machine-learning>, 2015.
- FULKERSON, B.; VEDALDI, A.; SOATTO, S. Class segmentation and object localization with superpixel neighborhoods. In: *2009 IEEE 12th International Conference on Computer Vision*. [S.l.: s.n.], 2009. p. 670–677.
- Ge, Z.; Chen, X. Dynamic probabilistic latent variable model for process data modeling and regression application. *IEEE Transactions on Control Systems Technology*, v. 27, n. 1, p. 323–331, 2019.
- GIRSHICK, R.; DONAHUE, J.; DARRELL, T.; MALIK, J. Region-based convolutional networks for accurate object detection and segmentation. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, v. 38, n. 1, p. 142–158, 2016.
- GOODFELLOW, I.; WARDE-FARLEY, D. Pylearn2: a machine learning research library. *arXiv preprint arXiv: 1401.0001*, p. 1–9, 01 2013.
- GOULD, S.; RODGERS, J.; COHEN, D.; ELIDAN, G.; KOLLER, D. Multi-class segmentation with relative location prior. In: . [S.l.: s.n.], 2008. v. 80, p. 300–316.
- GUYON, I.; SUN-HOSOYA, L.; BOULLÉ, M.; ESCALANTE, H. J.; ESCALERA, S.; LIU, Z.; JAJETIC, D.; RAY, B.; SAEED, M.; SEBAG, M.; STATNIKOV, A.; TU, W.-W.; VIEGAS, E. Analysis of the automl challenge series 2015–2018. In: _____. *Automated Machine Learning: Methods, Systems, Challenges*. Cham: Springer International Publishing, 2019. p. 177–219. ISBN 978-3-030-05318-5. Disponível em: <https://doi.org/10.1007/978-3-030-05318-5_10>.
- H, L.; K, S.; Y., Y. Darts: Differentiable architecture search. 2018.
- HAO, X.; ZHANG, G.; MA, S. Technical survey deep learning. *Int. J. Semantic Computing*, World Scientific, 2016.
- HE, X.; ZHAO, K.; CHU, X. Automl: A survey of the state-of-the-art. 2020.

HENRY, E.; HOFRICHTER, J. [8] singular value decomposition: Application to analysis of experimental data. *Methods in enzymology*, Elsevier, v. 210, p. 129–192, 1992.

HRENNIKOFF, A. Solution of problems of elasticity by the framework method. *J. Appl. Mech.*, 1941. Disponível em: <<https://ci.nii.ac.jp/naid/10015257315/en/>>.

HUTTER, F.; HOOS, H. H.; LEYTON-BROWN, K. Sequential model-based optimization for general algorithm configuration. In: COELLO, C. A. C. (Ed.). *Learning and Intelligent Optimization*. Berlin, Heidelberg: Springer Berlin Heidelberg, 2011. p. 507–523. ISBN 978-3-642-25566-3.

HUTTER, F.; KOTTHOFF, L.; VANSCHOREN, J. *Automated machine learning: methods, systems, challenges*. [S.l.]: Springer Nature, 2019.

INC., P. T. Collaborative data science. Plotly Technologies Inc., Montreal, QC, 2015. Disponível em: <<https://plot.ly>>.

JIA, Y.; SHELHAMER, E.; DONAHUE, J.; KARAYEV, S.; LONG, J.; GIRSHICK, R. B.; GUADARRAMA, S.; DARRELL, T. Caffe: Convolutional architecture for fast feature embedding. *CoRR*, abs/1408.5093, 2014. Disponível em: <<http://arxiv.org/abs/1408.5093>>.

JIN, H.; SONG, Q.; HU, X. Auto-keras: An efficient neural architecture search system. In: *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*. New York, NY, USA: Association for Computing Machinery, 2019. (KDD '19), p. 1946–1956. ISBN 9781450362016. Disponível em: <<https://doi.org/10.1145/3292500.3330648>>.

JOSHI, A. V. *Machine Learning and Artificial Intelligence*. [S.l.]: Springer, 2020.

KAHNG, M.; ANDREWS, P. Y.; KALRO, A.; CHAU, D. H. Activis: Visual exploration of industry-scale deep neural network models. *CoRR*, abs/1704.01942, 2017. Disponível em: <<http://arxiv.org/abs/1704.01942>>.

KAHNG, M.; ANDREWS, P. Y.; KALRO, A.; CHAU, D. H. Activis: Visual exploration of industry-scale deep neural network models. *IEEE Transactions on Visualization and Computer Graphics*, v. 24, n. 1, p. 88–97, 2018. ISSN 1077-2626, 2160-9306.

KOMER, B.; BERGSTRA, J.; ELIASMITH, C. Hyperopt-sklearn. In: _____. *Automated Machine Learning: Methods, Systems, Challenges*. Cham: Springer International Publishing, 2019. p. 97–111. ISBN 978-3-030-05318-5. Disponível em: <https://doi.org/10.1007/978-3-030-05318-5_5>.

KRAMER, M. A. Nonlinear principal component analysis using autoassociative neural networks. *AIChE Journal*, v. 37, n. 2, p. 233–243, 1991. Disponível em: <<https://aiche.onlinelibrary.wiley.com/doi/abs/10.1002/aic.690370209>>.

LECUN, Y.; BENGIO, Y.; HINTON, G. Deep learning. *nature*, v. 521, p. 436–444, 2015.

LEE, J. H.; SHIN, J.; REALFF, M. J. Machine learning: Overview of the recent progresses and implications for the process systems engineering field. *Computers and Chemical Engineering*, v. 114, p. 111–121, 2018.

LEI, Y.; YAO, X.; CHEN, W.; ZHANG, J.; MEHNEN, J.; YANG, E. Multiple object detection of workpieces based on fusion of deep learning and image processing*. In: *2020 International Joint Conference on Neural Networks (IJCNN)*. [S.l.: s.n.], 2020. p. 1–7.

LI, Y.; SUN, J.; TANG, C.-K.; SHUM, H.-Y. Lazy snapping. *ACM Trans. Graph.*, Association for Computing Machinery, New York, NY, USA, v. 23, n. 3, p. 303–308, ago. 2004. ISSN 0730-0301. Disponível em: <<https://doi.org/10.1145/1015706.1015719>>.

LONG, J.; SHELHAMER, E.; DARRELL, T. *Fully Convolutional Networks for Semantic Segmentation*. 2015.

LUNDBERG, S. M.; LEE, S.-I. A unified approach to interpreting model predictions. Curran Associates, Inc., p. 4765–4774, 2017.

MAATEN, L. Van der; HINTON, G. Visualizing data using t-sne. *Journal of machine learning research*, v. 9, n. 11, 2008.

MACQUEEN, J. et al. Some methods for classification and analysis of multivariate observations. In: OAKLAND, CA, USA. *Proceedings of the fifth Berkeley symposium on mathematical statistics and probability*. [S.l.], 1967. v. 1, n. 14, p. 281–297.

MANN, V.; VENKATASUBRAMANIAN, V. Predicting chemical reaction outcomes: A grammar ontology-based transformer framework. *AIChE Journal*, v. 67, n. 3, p. e17190, 2021. Disponível em: <<https://aiche.onlinelibrary.wiley.com/doi/abs/10.1002/aic.17190>>.

MCINNIS, L.; HEALY, J.; MELVILLE, J. Umap: uniform manifold approximation and projection for dimension reduction. 2020.

MIDWAY, S. R. Principles of effective data visualization. *Patterns*, v. 1, n. 9, p. 100141, 2020. ISSN 2666-3899. Disponível em: <<https://www.sciencedirect.com/science/article/pii/S2666389920301896>>.

Ming, Y.; Qu, H.; Bertini, E. Rulematrix: Visualizing and understanding classifiers with rules. *IEEE Transactions on Visualization and Computer Graphics*, v. 25, n. 1, p. 342–352, 2019.

MITCHELL, T. M. *Machine Learning*. 1. ed. USA: McGraw-Hill, Inc., 1997. ISBN 0070428077.

Mount III, E. M. Impact of metallizer process controls on optical and gas barrier uniformity. *Journal of plastic film and sheeting*, EMMOUNT Technologies, v. 24, n. 3-4, p. 203–211, 2008.

NING, C.; YOU, F. Optimization under uncertainty in the era of big data and deep learning: When machine learning meets mathematical programming. *Computers & Chemical Engineering*, v. 125, p. 434–448, 2019. ISSN 0098-1354. Disponível em: <<https://www.sciencedirect.com/science/article/pii/S0098135419300687>>.

NWANKPA, C.; IJOMAH, W.; GACHAGAN, A.; MARSHALL, S. Activation functions: Comparison of trends in practice and research for deep learning. *CoRR*, abs/1811.03378, 2018. Disponível em: <<http://arxiv.org/abs/1811.03378>>.

OLSON, R. S.; MOORE, J. H. Tpot: A tree-based pipeline optimization tool for automating machine learning. In: _____. *Automated Machine Learning: Methods, Systems, Challenges*. Cham: Springer International Publishing, 2019. p. 151–160. ISBN 978-3-030-05318-5. Disponível em: <https://doi.org/10.1007/978-3-030-05318-5_8>.

ONO, J. P.; CASTELO, S.; LOPEZ, R.; BERTINI, E.; FREIRE, J.; SILVA, C. Pipelineprofiler : A visual analytics tool for the exploration of automl pipelines. 2020. Disponível em: <arXiv:2005.00160v2>.

OSCO, L. P.; ARRUDA, M. d. S. de; GONÇALVES, D. N.; DIAS, A.; BATISTOTI, J.; SOUZA, M. de; GOMES, F. D. G.; RAMOS, A. P. M.; JORGE, L. A. d. C.; LIESENBERG, V. et al. A cnn approach to simultaneously count plants and detect plantation-rows from uav imagery. *arXiv preprint arXiv:2012.15827*, 2020.

OVERCASH, M. R. Perspective on advanced manufacturing and progress on improvement in societal well-being. *Journal of Advanced Manufacturing and Processing*, v. 1, n. 3, p. 2–3, 2019.

OZISIKYILMAZ, B.; MEMIK, G.; CHOUDHARY, A. Efficient system design space exploration using machine learning techniques. In: IEEE. *2008 45th ACM/IEEE Design Automation Conference*. [S.l.], 2008. p. 966–969.

PATEL, K.; BANCROFT, N.; DRUCKER, S. M.; FOGARTY, J.; KO, A. J.; LANDAY, J. Gestalt: Integrated support for implementation and analysis in machine learning. In: _____. New York, NY, USA: Association for Computing Machinery, 2010. p. 37–46. ISBN 9781450302715. Disponível em: <<https://doi.org/10.1145/1866029.1866038>>.

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.; DUCHESNAY Édouard. Scikit-learn: Machine learning in python. *Journal of Machine Learning Research*, v. 12, n. 85, p. 2825–2830, 2011. Disponível em: <<http://jmlr.org/papers/v12/pedregosa11a.html>>.

PERRY, M.; LENTZ, R. Susceptors in microwave packaging. *Development of Packaging and Products for Use in Microwave Ovens*, Woodhead Publishing, p. 207–236, 2009.

RADHAKRISHNAN, R. A survey of multiscale modeling: Foundations, historical milestones, current status, and future prospects. *AIChE Journal*, v. 67, n. 3, p. e17026, 2021. Disponível em: <<https://aiche.onlinelibrary.wiley.com/doi/abs/10.1002/aic.17026>>.

RASMUSSEN, C. E.; WILLIAMS, C. K. I. *Gaussian Processes for Machine Learning (Adaptive Computation and Machine Learning)*. [S.l.]: The MIT Press, 2005. ISBN 026218253X.

REDMON, J.; FARHADI, A. *YOLOv3: An Incremental Improvement*. 2018.

REN, D.; AMERSHI, S.; LEE, B.; SUH, J.; WILLIAMS, J. D. Squares: Supporting interactive performance analysis for multiclass classifiers. *IEEE Transactions on Visualization and Computer Graphics*, v. 23, n. 1, p. 61–70, 2017. ISSN 1077-2626.

REYNOLDS, D. A. Gaussian mixture models. *Encyclopedia of biometrics*, Springer City, v. 741, p. 659–663, 2009.

SHAHRIARI, B.; SWERSKY, K.; WANG, Z.; ADAMS, R. P.; FREITAS, N. de. Taking the human out of the loop: A review of bayesian optimization. *Proceedings of the IEEE*, v. 104, n. 1, p. 148–175, 2016.

SHIN, J.; BADGWELL, T. A.; LIU, K.-H.; LEE, J. H. Reinforcement learning –overview of recent progress and implications for process control. *Computers and Chemical Engineering*, Elsevier, v. 127, n. 3-4, p. 282–294, 2019.

- SHIN, J.; LEE, J. H.; REALFF, M. J. Operational planning and optimal sizing of microgrid considering multi-scale wind uncertainty. *Applied Energy*, v. 195, p. 616 – 633, 2017.
- SUTTON, R. S.; BARTO, A. G. *Reinforcement Learning: An Introduction*. Second. The MIT Press, 2018. Disponível em: <<http://incompleteideas.net/book/the-book-2nd.html>>.
- SUTTON, R. S.; BARTO, A. G. *Reinforcement Learning An introduction*. 2. ed. [S.l.: s.n.], 2018. ISBN 978-0-262-19398-6.
- THORNTON, C.; HUTTER, F.; HOOS, H. H.; LEYTON-BROWN, K. Auto-weka: Automated selection and hyper-parameter optimization of classification algorithms. *CoRR*, abs/1208.3719, 2012. Disponível em: <<http://arxiv.org/abs/1208.3719>>.
- TOPMET, A. *TopMet 2450 Operating Instructions*. [S.l.: s.n.], 2016.
- VENKATASUBRAMANIAN, V. The promise of artificial intelligence in chemical engineering: Is it here, finally? *AIChE Journal*, v. 65, n. 2, p. 466–478, 2019. Disponível em: <<https://aiche.onlinelibrary.wiley.com/doi/abs/10.1002/aic.16489>>.
- WANG, J.; LIU, B.; XU, K. Semantic segmentation of high-resolution images. *Sci. China Inf. Sci.*, v. 60, n. 3, p. 101–123, 2017.
- WEI, S.; ZHAO, X.; MIAO, C. A comprehensive exploration to the machine learning techniques for diabetes identification. In: IEEE. *2018 IEEE 4th World Forum on Internet of Things (WF-IoT)*. [S.l.], 2018. p. 291–295.
- WEICHERT, D.; LINK, P.; STOLL, A.; RÜPING, S.; IHLENFELDT, S.; WROBEL, S. A review of machine learning for the optimization of production processes. *The International Journal of Advanced Manufacturing Technology*, Springer, v. 104, n. 5, p. 1889–1902, 2019.
- WENG, Z. From conventional machine learning to automl. *IOP Conf. Series: Journal of Physics: Conf. Series*, IOP Publishing, v. 1207, 2019.
- Wongsuphasawat, K.; Smilkov, D.; Wexler, J.; Wilson, J.; Mané, D.; Fritz, D.; Krishnan, D.; Viégas, F. B.; Wattenberg, M. Visualizing dataflow graphs of deep learning models in tensorflow. *IEEE Transactions on Visualization and Computer Graphics*, v. 24, n. 1, p. 1–12, 2018.
- WORTMANN, F.; FLÜCHTER, K. Internet of things. *Business & Information Systems Engineering*, Springer, v. 57, n. 3, p. 221–224, 2015.
- XU, P.; MEI, H.; REN, L.; CHEN, W. Vidx: Visual diagnostics of assembly line performance in smart factories. *IEEE Transactions on Visualization and Computer Graphics*, v. 23, n. 1, p. 291–300, 2017. ISSN 1077-2626.
- YANG, C.; FAN, J.; WU, Z.; UDELL, M. Efficient automl pipeline search with matrix and tensor factorization. 2020.
- YANG, Y.; HALLMAN, S.; RAMANAN, D.; FOWLKES, C. Layered object detection for multi-class segmentation. In: *2010 IEEE Computer Society Conference on Computer Vision and Pattern Recognition*. [S.l.: s.n.], 2010. p. 3113–3120.
- YUE, H.; QIN, S.; MARKLE, R.; NAUERT, C.; GATTO, M. Fault detection of plasma etchers using optical emission spectra. *IEEE Transactions on Semiconductor Manufacturing*, v. 13, n. 3, p. 374–385, 2000.

ZHANG, K.; ZHANG, Y.; WANG, M. A unified approach to interpreting model predictions scott. *Neural Information Processing Systems*, v. 16, n. 3, p. 426–430, 2012.

ZHOU, F.; LIN, X.; LUO, X.; ZHAO, Y.; CHEN, Y.; CHEN, N.; GUI, W. Visually enhanced situation awareness for complex manufacturing facility monitoring in smart factories. *Journal of Visual Languages & Computing*, v. 44, p. 58–69, 2018. ISSN 1045-926X. Disponível em: <<https://www.sciencedirect.com/science/article/pii/S1045926X17301829>>.

ZHOU, X.; ZHUO, J.; KRÄHENBÜHL, P. *Bottom-up Object Detection by Grouping Extreme and Center Points*. 2019.

ZIMMER, L.; LINDAUER, M.; HUTTER, F. *Auto-PyTorch Tabular: Multi-Fidelity MetaLearning for Efficient and Robust AutoDL*. 2020.

ZITNICK, C.; KANG, S. Stereo for image-based rendering using image over-segmentation. In: . [S.l.: s.n.], 2008. v. 75, p. 49–65.