

UNIVERSIDADE DE SÃO PAULO

Instituto de Ciências Matemáticas e de Computação

**Deteccção de anomalias em sensores de poços
submarinos com uso de redes neurais artificiais**

Gustavo Garcia Momm

Monografia - MBA em Inteligência Artificial e Big Data

SERVIÇO DE PÓS-GRADUAÇÃO DO ICMC-USP

Data de Depósito:

Assinatura: _____

Gustavo Garcia Momm

Detecção de anomalias em sensores de poços submarinos com uso de redes neurais artificiais

Monografia apresentada ao Departamento de Ciências de Computação do Instituto de Ciências Matemáticas e de Computação, Universidade de São Paulo - ICMC/USP, como parte dos requisitos para obtenção do título de Especialista em Inteligência Artificial e Big Data.

Área de concentração: Inteligência Artificial

Orientadora: Prof. Dr. Anna Helena Reali Costa

Versão original

São Carlos

2022

Ficha catalográfica elaborada pela Biblioteca Prof. Achille Bassi
e Seção Técnica de Informática, ICMC/USP,
com os dados inseridos pelo(a) autor(a)

G216d Garcia Momm, Gustavo
 Detecção de anomalias em sensores de poços
 submarinos com uso de redes neurais artificiais /
 Gustavo Garcia Momm; orientador Anna Helena Reali
 Costa. -- São Carlos, 2022.
 48 p.

 Trabalho de conclusão de curso (MBA em
 Inteligência Artificial e Big Data) -- Instituto de
 Ciências Matemáticas e de Computação, Universidade
 de São Paulo, 2022.

 1. Detecção de Anomalias. 2. Redes Neurais
 Artificiais. 3. Multi-Layer Perceptron. I. Reali
 Costa, Anna Helena, orient. II. Título.

Gustavo Garcia Momm

Anomaly detection in subsea well sensors using artificial neural networks

Monograph presented to the Departamento de Ciências de Computação do Instituto de Ciências Matemáticas e de Computação, Universidade de São Paulo - ICMC/USP, as part of the requirements for obtaining the title of Specialist in Artificial Intelligence and Big Data.

Concentration area: Artificial Intelligence

Advisor: Prof. Anna Helena Reali Costa

Original version

São Carlos

2022

RESUMO

Momm, G.G. **Detecção de anomalias em sensores de poços submarinos com uso de redes neurais artificiais**. 2022. 48p. Monografia (MBA em Inteligência Artificial e Big Data) - Instituto de Ciências Matemáticas e de Computação, Universidade de São Paulo, São Carlos, 2022.

No contexto geral da indústria, há uma demanda crescente pelo aumento da segurança operacional, produtividade, qualidade e eficiência energética. A indústria de óleo e gás se enquadra nesse contexto e tem aplicado diferentes técnicas computacionais com o intuito, principalmente, de aprimorar a segurança operacional. Essas diferentes técnicas utilizam dados obtidos a partir de sensores de pressão e temperatura instalados no poço e na plataforma de produção ao qual o poço está interligado. Os dados obtidos por sensores estão sujeitos a erros intrínsecos desses equipamentos ou dos sistemas aos quais estão interligados, podendo comprometer análises dependentes desses dados. Esse trabalho investiga dois modelos de aprendizado de máquina supervisionado utilizando redes neurais artificiais para detectar dados anômalos de sensores de poços submarinos: o modelo de classificação e o modelo de previsão. Em ambos modelos foram utilizadas redes neurais artificiais do tipo *Multi-Layer Perceptron*. Dados reais de um poço submarino da costa brasileira foram usados para as análises. O modelo de classificação utilizou dados rotulados para treinamento e validação da rede neural artificial e o modelo de previsão utilizou a técnica *Sliding-Window* para prever dados a partir de trechos anteriores das séries temporais. O modelo de previsão apresentou melhores resultados quando comparado com o modelo de classificação e tem maior potencial de aplicação prática, pois permite a identificação de um dado anômalo medido com base na comparação com o dado predito e um critério de desvio previamente definido.

Palavras-chave: Detecção de Anomalias, Redes Neurais Artificiais, Multi-Layer Perceptron.

ABSTRACT

Momm, G.G **Anomaly detection in subsea well sensors using artificial neural networks**. 2022. 48p. Monograph (MBA in Artificial Intelligence and Big Data) - Instituto de Ciências Matemáticas e de Computação, Universidade de São Paulo, São Carlos, 2022.

In the general context of the industry, there is a growing demand for increased operational safety, productivity, quality and energy efficiency. The oil and gas industry fits into this context and has applied different computational techniques in order, mainly, to improve operational safety. These different techniques use data obtained from pressure and temperature sensors installed in the well and on the production platform to which the well is connected. The data obtained by sensors are subject to intrinsic errors of these equipment or of the systems to which they are interconnected, which may compromise analyzes dependent on these data. This work investigates two supervised machine learning models using artificial neural networks to detect anomalous data from subsea well sensors: a classification model and a prediction model. In both models, artificial neural networks of the *Multi-Layer Perceptron* type were used. Real data from a subsea well of the Brazilian coast were used for the analyses. The classification model used labeled data for training and validation of the artificial neural network and the prediction model used the *Sliding-Window* technique to predict data from preceding slice of the time series. The prediction model presented better results when compared to the classification model and has greater potential for practical application, as it allows the identification of an anomalous data based on the comparison with the predicted data and a previously defined deviation criterion.

Keywords: Anomaly Detection, Artificial Neural Network, Multi-Layer Perceptron.

LISTA DE ABREVIATURAS E SIGLAS

FN	<i>False Negative</i>
FP	<i>False Positive</i>
LSTM	<i>Long Short-Term Memory</i>
MAPE	<i>Mean Absolute Percentage Error</i>
MLP	<i>Multi-Layer Perceptron</i>
MSE	<i>Mean Squared Error</i>
PCA	<i>Principal Component Analysis</i>
PCK	<i>Production Choke</i>
PDG	<i>Permanent Downhole Gauge</i>
RMSE	<i>Root-Mean-Square Error</i>
RNA	Rede Neural Artificial
TDNN	<i>Time Delay Neural Network</i>
TP	<i>True Positive</i>
TPT	<i>Temperture and Pressure Transmitter</i>
USP	Universidade de São Paulo
USPSC	Campus USP de São Carlos

SUMÁRIO

1	INTRODUÇÃO	15
1.1	Justificativa e Motivação	15
1.2	Questões de Pesquisa e Objetivos	17
1.3	Organização do Texto	17
2	FUNDAMENTAÇÃO TEÓRICA	19
2.1	Detecção de Anomalias	19
2.2	Redes Neurais Artificiais	20
2.3	Estado da arte	23
3	BASE DE DADOS UTILIZADA	27
3.1	Base de dados	27
3.2	Análise da qualidade dos dados	28
3.3	Análise da correlação entre variáveis	31
3.4	Análise estatística dos dados	33
3.5	Preparação dos dados	33
4	DESENVOLVIMENTO E RESULTADOS	37
4.1	Modelo de Classificação	37
4.2	Modelo de Previsão	39
4.3	Avaliação dos Resultados	41
4.3.1	Modelo de Classificação	41
4.3.2	Modelo de Previsão	42
5	CONCLUSÃO	45
	REFERÊNCIAS	47

1 INTRODUÇÃO

No contexto geral da indústria, há uma demanda crescente pelo aumento da segurança operacional, produtividade, qualidade e eficiência energética (1). A indústria de óleo e gás se enquadra nesse contexto, sendo motivada principalmente por aspectos ambientais, regulatório e econômicos.

A complexidade envolvida no processo de extração de óleo e gás, desde a fase inicial de exploração até a final com a distribuição, exige bastante cuidado na execução dos processos realizados ao longo de toda a cadeia produtiva do petróleo (2). O acidente ocorrido com o navio de perfuração DeepWater Horizon durante a perfuração do poço de Macondo em 2010 é um exemplo do potencial de impacto que falhas de equipamentos de segurança podem gerar. Nesse acidente 11 pessoas morreram e ele ficou marcado como o maior acidente ambiental da história dos Estados Unidos (1).

Na busca pelo aumento da segurança operacional de poços de petróleo diferentes técnicas computacionais têm sido aplicadas (1), como algoritmos de aprendizagem profunda para detectar anomalias de produção (2) e *Digital Twin* para monitoramento da integridade de poços em tempo real (3). Essas diferentes técnicas utilizam dados obtidos a partir de sensores de pressão e/ou temperatura instalados no poço e na plataforma de produção ao qual o poço está interligado. A disponibilização em terra de dados de sensores instalados em ambientes remotos pode ser comprometida caso esses sensores, ou o sistema de transmissão dos dados, apresentem falhas.

Na figura 1 são ilustrados os principais sensores instalados em poços de petróleo submarino que permitem a medição das seguintes grandezas: pressão e temperatura no sensor do fundo do poço (PDG), pressão e temperatura na cabeça do poço (TPT) e pressão e temperatura na plataforma (PCK).

1.1 Justificativa e Motivação

Dados obtidos por sensores estão sujeitos a erros intrínsecos desses equipamentos ou dos sistemas aos quais estão interligados. Parte desses erros é conhecida desde sua concepção e parte se revela durante sua vida útil. Rotinas de manutenção e calibração visam manter esses erros dentro de uma faixa prevista em projeto. Sensores instalados em locais remotos, como poços de petróleo, não permitem que essas rotinas sejam executadas frequentemente, tanto por razões de segurança como econômicas.

Na figura 2 são apresentadas as categorias de falhas de sensores (4): atraso, deslocamento e alteração de valores. Soma-se a essas categorias a indisponibilidade de dados. No caso do banco de dados 3W (1), que compreende dados reais de poços de petróleo da

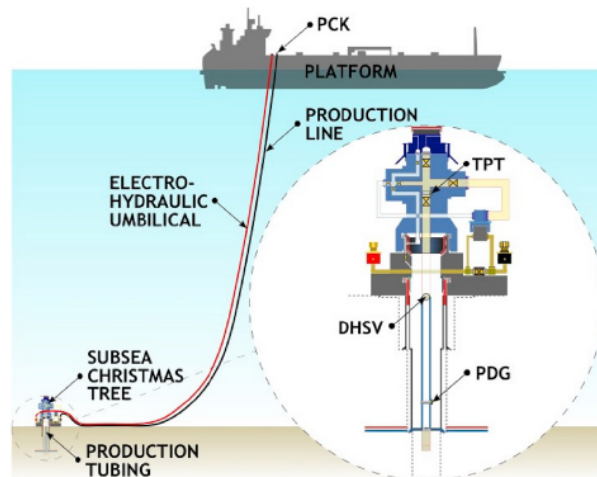


Figura 1 – Esquemático de um poço submarino com localização dos sensores (1).

costa brasileira, os dados considerados ausentes representam 31,17% do banco de dados.

Delay	Offset			Stuck-at
	sporadic	permanent	stochastic	
constant y t ①	outlier y t ②	constant y t ③	constant y expected real t ④	at zero y t ⑤
variable y t ⑥	spike y t ⑦	value correlated y t ⑧	value correlated y expected real t ⑨	at X y t ⑩
omissions / broken link y omissions t ⑪		time correlated y t ⑫	time correlated y expected real t ⑬	saturation y t ⑭

Figura 2 – Categoria de falhas de sensores (4). Linha tracejada representa o valor real e a linha sólida representa a medida com erro.

A identificação de comportamento anômalo em sensores é possível por diferentes técnicas. Teh et al (5) realizaram uma extensa pesquisa bibliográfica e identificaram que cerca de 40% dos trabalhos avaliados utilizaram as técnicas de Análise de Componentes Principais (PCA) e Redes Neurais Artificiais (RNA) para detecção de erros em sensores. As RNA apresentam vantagens frente aos métodos baseados em regras e possuem generalização melhor, sendo portanto mais adequadas ao problema proposto (4).

Nesse trabalho são utilizadas redes neurais artificiais para identificar dados anômalos de um dos sensores de poços submarinos a partir da dados históricos do conjunto de sensores.

1.2 Questões de Pesquisa e Objetivos

Esse trabalho tem o objetivo de avaliar modelos de aprendizagem de máquina baseados em redes neurais artificiais (RNA) para detectar dados anômalos de sensores de poços submarinos, podendo esse problema ser desdobrado em duas questões a serem abordadas:

1. Modelos de aprendizagem de máquina baseados em redes neurais artificiais permitem a identificação de estados anômalos de sensores de poços submarinos de petróleo?
2. Os dados coletados no histórico do poço são suficientes para o treinamento de um modelo RNA de forma que o erro observado esteja similar aos encontrados na literatura?

1.3 Organização do Texto

O texto está organizado da seguinte forma. No capítulo 2 é apresentada a fundamentação teórica do trabalho, com uma breve descrição da detecção de anomalias e das redes neurais artificiais. A fundamentação é complementada com a apresentação de trabalhos relevantes para esse desenvolvimento.

No capítulo 4 é apresentado o banco de dados selecionado, suas principais características e o conjunto de dados a ser utilizados nas análises. Esse dados passam por uma preparação e são separados três conjuntos de teste na seção 3.5. As seções 4.1 e 4.2 são dedicadas para a descrição dos modelos utilizados e suas métricas de desempenho. Os resultados de cada modelo são apresentados na seção 4.3.

Por fim, no capítulo 5 são consolidadas as conclusões e apresentada uma proposta de trabalho futuro.

2 FUNDAMENTAÇÃO TEÓRICA

Neste capítulo são abordadas técnicas de detecção de anomalias e uma breve descrição do conceito de Redes Neurais Artificiais. O capítulo é complementado com um panorama sobre os trabalhos relevantes para o presente estudo.

2.1 Detecção de Anomalias

Detecção de *outliers*, também conhecida como detecção de anomalia em algumas literaturas, é, há muito tempo, um tema de pesquisa importante nas áreas de mineração de dados e estatística. O principal objetivo da detecção de *outliers* é a identificação de dados que são claramente diferentes ou inconsistentes com os demais dados do conjunto (6). A detecção de anomalias, no contexto desse trabalho, consiste na identificação de desvios a partir de um comportamento esperado para uma variável.

Os métodos empregados para essa detecção variam desde a comparação de dados em relação a limites estabelecidos até algoritmos complexos de filtragem de sinais (4). Na figura 3 é ilustrado o processo de identificação de anomalias, na qual dados históricos são utilizados para auxiliar na determinação de condições classificadas como normais e anômalas (4).

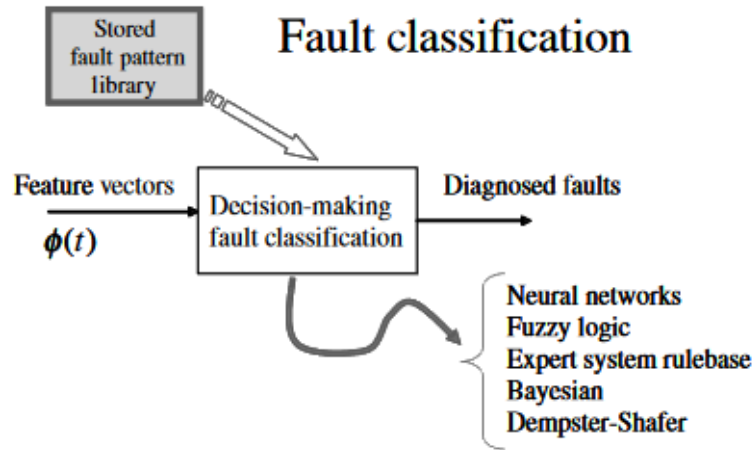


Figura 3 – Visão geral de um processo de detecção de anomalias (7).

Hodge e Austin (8) apresentam três categorias de métodos de detecção de anomalias, baseadas na filosofia do modelo de detecção:

- Tipo 1: As anomalias (*outliers*) são determinadas sem conhecimento prévio dos dados. É essencialmente uma abordagem similar à classificação não supervisionada. Essa abordagem analisa a distribuição estatística dos dados destacando os pontos remotos e identificando-os como potenciais *outliers*;

- Tipo 2: São modeladas a normalidade e a anormalidade. Essa abordagem é análoga à classificação supervisionada e requer dados previamente classificados entre normais e anormais;
- Tipo 3: São modeladas somente normalidades ou em poucos casos a anormalidade. Autores geralmente identificam essa técnica como detecção de novidade ou reconhecimento de novidade. É análoga às tarefas de detecção ou reconhecimento semi-supervisionados, pois a classe normal é fornecida e o algoritmo aprende a reconhecer anormalidades. Essa abordagem requer dados pré-classificados, mas somente aprende com dados marcados como normais.

Os autores de (9) consolidaram em um estudo técnicas para detecção de anomalias em dados temporais. Foi adotada uma divisão focada na característica dos dados e nos cenários de aplicação, conforme figura 4. Para cada uma das abordagens são identificadas técnicas de detecção de anomalia que se enquadram na divisão apresentada por Hodge e Austin (8).

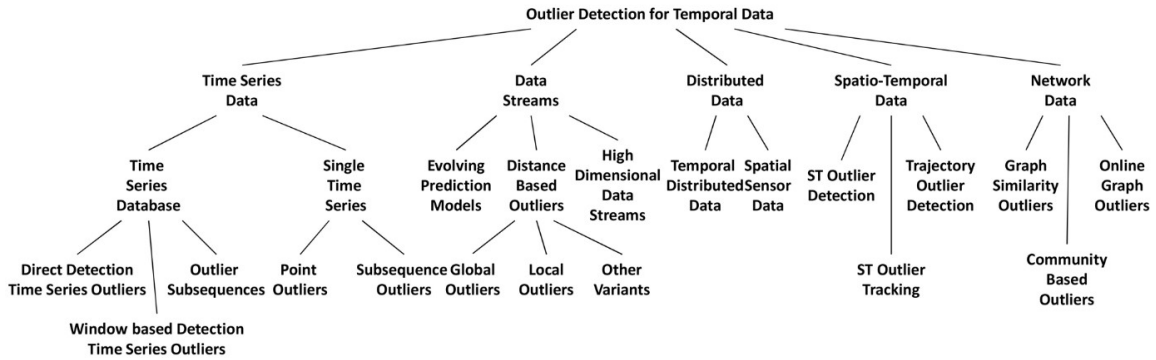


Figura 4 – Abordagens para detecção de anomalias em dados temporais conforme tipo de dados (9).

As redes neurais artificiais são adotadas em diferentes modelos de detecção, destacando a aplicação em modelos de classificação supervisionada, como em Jäger *et al* (4), e aplicação em modelos de previsão, nos quais os valores medidos são comparados com os valores preditos para identificar possíveis anomalias, como em Saad *et al* (10) e Yu *et al* (6).

2.2 Redes Neurais Artificiais

As redes neurais artificiais (RNA) são algoritmos de aprendizagem de máquina formulados a partir do modelo não linear dos neurônios, denominado *perceptron*. De forma geral, uma RNA recebe um sinal de entrada $\mathbf{x} = (x_1, x_2, \dots, x_m)$ e realiza operação com seus respectivos pesos w_{ki} . Esses sinais de entrada são combinados e a amplitude é ajustada com o uso de um *bias* b_k

$$v_k = \sum_{i=1}^m (w_{ki} \cdot x_i + b_k). \quad (2.1)$$

Por fim o sinal passa por uma função de ativação não linear, gerando o sinal de saída y_k (10),

$$y_k = \varphi(v_k). \quad (2.2)$$

Na figura 5 é ilustrado graficamente este processo.

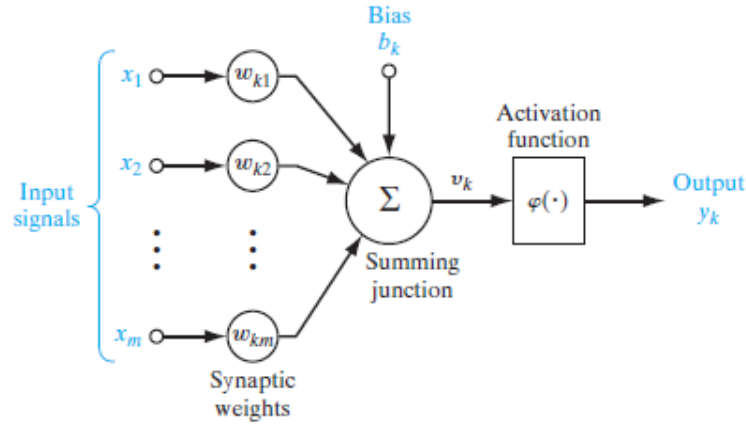


Figura 5 – Modelo não linear de um neurônio k (11).

As funções de ativação mais comuns são sigmóide,

$$\varphi(z) = \frac{1}{1 + e^{-z}}, \quad (2.3)$$

e ReLU, que é uma abreviação para unidade linear retificada,

$$\varphi(z) = \max(0, z). \quad (2.4)$$

Com o intuito de aumentar a capacidade da RNA de trabalhar estatísticas de mais alta ordem, são adicionadas camadas intermediárias (escondidas) a essas redes, de forma que é gerado um outro conjunto de conexões e uma dimensão adicional de conexões (11). Essas redes com múltiplas camadas com alimentação à frente, ou seja uma generalização do *perceptron*, são conhecidas como *Multi-Layer Perceptron* (MLP). Na figura 6 é ilustrada uma configuração de uma rede MLP.

Dado um conjunto de dados $D = \{(\mathbf{x}^{(i)}, y^{(i)})\}_{i=1}^{N_D}$ composto por N_D amostras e com entradas $\mathbf{x}^{(i)}$ e respectivas saídas $y^{(i)}$ conhecidas, o processo de treinamento de uma RNA consiste em encontrar os pesos $\mathbf{W} = \{w\}$ que minimizam uma função de perda (*loss function*), $L(\mathbf{W})$, tipicamente dada pelo Erro Quadrático Médio (MSE),

$$L(\mathbf{W}) = \frac{1}{N_D} \sum_{i=1}^{N_D} \left[\frac{1}{K} \sum_{k=1}^K \left(y_k^{(i)} - \varphi(v_k)^{(i)} \right)^2 \right], \quad (2.5)$$

sendo K o número de neurônios na camada de saída e v_k dado pela equação 2.1.

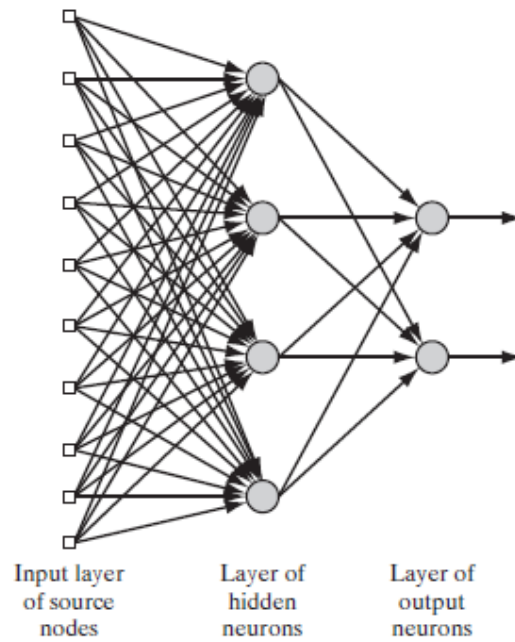


Figura 6 – Rede com uma camada escondida totalmente conectada (11).

O treinamento é realizado utilizando o algoritmo *back-propagation* para aumentar a precisão de uma MLP. O algoritmo *back-propagation* utiliza o erro associado a diferença entre a saída da MLP e a saída esperada para realizar ajustes nos pesos de cada camada da rede, começando pela camada da saída para a de entrada. Esse algoritmo emprega diferentes otimizadores para redução da função perda, como *gradient descent*, *adaptive gradient*, Adam, entre outros (10). A taxa com a qual os pesos são alterados para se aproximar do ótimo é chamada taxa de aprendizagem. Quanto menor a taxa de aprendizagem menor são as alterações nos pesos de uma iteração para outra. Por outro lado, se a taxa de aprendizagem for muito grande para acelerar o treinamento, as grandes alterações nos pesos podem tornar a rede instável (11).

No treinamento de redes neurais é utilizado o método de *batch*, no qual ajustes nos pesos são realizados após a apresentação de um grupo de exemplos (*batch*) do conjunto de treinamento que constituem uma época de treinamento (11). O tamanho dos *batches* e número de épocas são parâmetros os principais parâmetros para o treinamento de redes neurais.

Em geral, o algoritmo *back-propagation* pode não convergir e não há critérios bem definidos para a operação. Porém, há alguns critérios razoáveis que podem ser utilizados para interromper os ajustes dos pesos (11). Na prática é considerado que houve convergência quando a taxa de variação do erro por época é suficientemente pequeno. Haykin (11) apresenta que a taxa de variação do erro quadrático médio é normalmente considerada pequena quando está na faixa de 0,1 a 1% por época.

2.3 Estado da arte

Um desenvolvimento relevante no contexto desse trabalho foi publicado por Guo e Nurre (12), no qual os pesquisadores aplicaram RNA para detecção de falhas de sensores em motores principais de *Space Shuttle*. Mais de uma centena de sensores equipam esses sistemas, porém os autores definiram 10 sensores críticos para a operação normal. Além de identificar a medição de sensor que não está consistente com as demais (detecção de falhas), os autores propuseram um método para reconstruir o sinal do sensor em falha. Os autores trataram as falhas de sensores individualmente, não sendo, portanto, consideradas falhas simultâneas. Foi empregada uma rede MLP para a detecção de falhas com 10 neurônios na camada de entrada, duas camadas escondidas com 30 neurônios cada e uma camada de saída com 10 neurônios representando o nível de confiança de cada sensor. O método empregado seguiu os seguintes passos:

1. Selecionar aleatoriamente 1 dos 150 conjuntos de medidas;
2. Selecionar aleatoriamente 1 dos 10 sensores a ser treinado;
3. Gerar um ruído Gaussiano aleatório ω com média zero e desvio padrão $\sigma = 1.5\epsilon_i$, onde $\pm\epsilon_i$ é o range de medidas válidas para o sensor no conjunto de medições S_i . A adição de ruído ω gera o conjunto S_i^* . Essa seleção de ruído gera aproximadamente 50% de medidas fora do intervalo das amostras de treinamento;
4. Se S_i^* está no intervalo de S_i então é definida a saída desejada da rede O_i como 0,9, caso contrário 0,1. Também são definidas todas as outras saídas como 0,9.
5. Ajustar os pesos de acordo com o algoritmo de *back propagation*;
6. Repetir os passos (1) a (5) até que o critério de parada.

Como resultado, os autores obtiveram acurácia acima de 90% para identificar falhas de sensor em um conjunto de teste contendo um milhão de amostras.

Teh *et al.* (5) realizaram uma revisão bibliográfica sobre publicações relacionadas a qualidade de dados de sensores, ao qual a detecção de falhas, ou anomalias, está abordado. Foram listados seis publicações que utilizaram RNA para detecção de anomalias, dessas destaca-se aqui o trabalho de Jäger *et al.* (4), que utilizou *Time-delay Neural Network* (TDNN) para detectar quatro tipos de anomalias: outliers, offset, ruído e congelamento em zero. TDNN é um tipo de RNA multi camadas *feed-forward* que permite o mapeamento entre valores passados e presentes pela análise de uma janela deslizante do sinal. A diferença entre a TDNN e o clássico *multilayer perceptron* (MLP) é que os neurônios não recebem somente a saída do neurônio anterior, mas também a saída com atraso (passado) desses neurônios (5). Os resultados indicaram precisão e revocação (*recall*) para cada uma das anomalias previstas conforme a Tabela 1. A RNA empregada pode detectar duas anomalias

com grande precisão e revocação, congelamento em zero e *offset* constante, com valores próximos ou iguais a 1, e falhou ao detectar *outlier* e ruído constante, uma vez que apresentou valores relativamente baixos para as métricas selecionadas.

Tabela 1 – Resultados de Precisão e Revocação obtidos por Jäger *et al.* (4)

Anomalia	Precisão	Revocação
Outlier	0,1	0,02
Offset Constante	1	1
Ruído Constante	0,47	0,40
Congelamento em zero	0,98	0,99

Na área de óleo e gás, Aggrey e Davies (13) utilizaram MLP para identificar anomalias em sensor de fundo de poço (PDG na figura 1). A abordagem adotada utilizou RNA para prever o valor de um sensor a partir de dados de outros 6 sensores e uma RNA com 2 camadas escondidas foi empregada. O dado predito foi então comparado com o dado medido para identificação de eventuais desvios. Foram geradas perturbações nos sinais dos sensores para avaliar a capacidade de identificação de condições anômalas, como os pontos A e B na figura 7. No ponto A a RNA recebeu dados corretos e estimou o valor corretamente, sendo os picos do sinal do sensor referentes a um comportamento anômalo desse sinal. No ponto B a RNA recebeu uma das entradas com erro associado e essa condição foi refletida no valor predito.

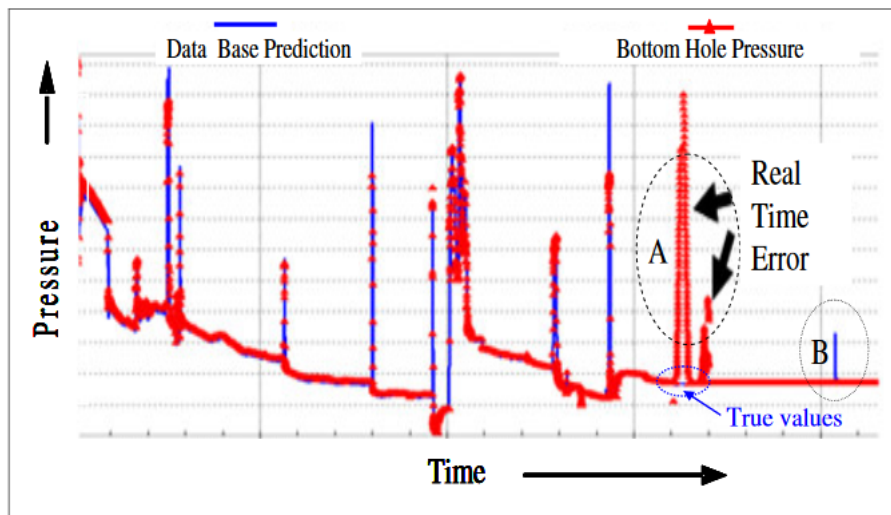


Figura 7 – Comparação entre valores preditos (linha azul) e valores medidos (linha vermelha). No ponto A a RNA teve como entrada dados corretos para prever valores de um sensor em falha e no ponto B a RNA teve uma de suas entradas com valor com erro (13).

Sobrinho *et al.* (2) utilizaram uma rede neural recorrente *Long Short-Term Memory* (LSTM) para identificar condições anômalas de poços de petróleo. Para isso foram utilizados dados de sensores disponíveis e de estados de válvulas (aberta, fechada ou percentual de abertura) para determinar as condições normais e anômalas. No primeiro momento

foram realizados estudos, através de árvore de decisão, para determinar comportamentos e correlações a partir dos dados de sensores para as combinações mais comuns de estados de válvulas. A resposta dessa análise foi aplicada ao treinamento da LSTM.

Saad *et al.* (10) também empregaram LSTM para detecção de falhas de linhas de ancoragem de plataformas de produção de petróleo. Nesse trabalho redes LSTM e MLP foram treinadas com dados de movimento da plataforma com o objetivo de prever o movimento horizontal, considerando que todas as linhas de ancoragem estavam íntegras. A falha de uma ou mais linhas de ancoragem de uma plataformas de produção de petróleo leva a uma resposta dinâmica, frente às condições ambientais (ondas, correntes e ventos), diferente da situação em que todas as linhas estão íntegras. Isso gera alteração na série temporal dos movimentos da plataforma. Essa falha foi simulada e a diferença entre os valores preditos pelas RNA (treinada sem falhas no sistema de amarração) e os valores medidos permitiram identificar a falha tanto para a LSTM como para o MLP.

Na tabela 2 são apresentados os modelos e estratégias utilizadas nos trabalhos analisados, bem como a abordagem proposta. Desses trabalhos conclui-se que RNA possuem grande capacidade de detectar anomalias mesmo em configurações mais simples.

Nesse trabalho são exploradas as redes MLP em dois modelos distintos para responder às duas questões apresentadas no item 1.2, sendo um modelo baseado na estratégia de classificação, no qual as observações utilizadas no treinamento e validação são classificadas em normais e anômalas, e um modelo de previsão, no qual dados de intervalos de tempo são utilizados para prever valores de uma variável em períodos posteriores.

Tabela 2 – Comparativo entre trabalhos analisados e abordagem proposta

Trabalho	Ano	RNA	Estratégia
Guo e Nurre (12)	1991	MLP	Aprendizado Supervisionado e Classificação
Jäger et al. (4)	2014	TDNN	Aprendizado Supervisionado e Classificação
Aggrey e Davies (13)	2007	MLP	Aprendizado Supervisionado e Previsão
Sobrinho <i>et al.</i> (2)	2020	LSTM	Aprendizado Supervisionado e Classificação
Saad <i>et al.</i> (10)	2021	MLP	Aprendizado Supervisionado e Previsão
		LSTM	Aprendizado Supervisionado e Previsão
Abordagem Proposta		MLP	Aprendizado Supervisionado e Classificação
		MLP	Aprendizado Supervisionado e Previsão

3 BASE DE DADOS UTILIZADA

No desenvolvimento desse trabalho é utilizado o banco de dados 3W (1), que apresenta conjuntos de dados reais de poços da costa brasileira.

Na seção 3.1 são apresentadas as informações presentes nesses conjuntos de dados. Na seção 3.2 é realizada uma análise da qualidade dos dados com o objetivo de se identificar se os presentes atendem ao objetivo proposto, bem como, na seção 3.3, é analisada a interdependência das variáveis, utilizando matriz de correlação. A análise de componentes principais permite extrair as variáveis mínimas a serem utilizadas. Finalmente, na seção 3.4, é analisado o comportamento dos dados no decorrer do tempo e são obtidos dados estatísticos básicos dos mesmos.

3.1 Base de dados

A base de dados 3W (1)¹ apresenta dados de 16 poços com frequência de observação constante de 1 Hz e são incorporados dados de pressão no fundo do poço (P-PDG), pressão e temperatura na cabeça de poço (P-TPT e T-TPT, respectivamente), pressão a montante da válvula de controle de fluxo na plataforma (P-MONT-CKP), temperatura a jusante da válvula de controle de fluxo na plataforma (T-JUS-CKP), pressão e temperatura a jusante da válvula de controle de fluxo de gás para elevação artificial (P-JUS-CKGL e T-JUS-CKGL, respectivamente) e vazão desse gás (QGL). A localização desses sensores é ilustrada na figura 1.

Os conjuntos de dados possuem tamanhos diferentes, conforme é mostrado na figura 8, sendo considerados relevantes para esse trabalho os conjuntos com mais de 1.10^6 observações, ou seja, os conjuntos rotulados nesse trabalho como w1, w2, w5, w6, w8 e w10.

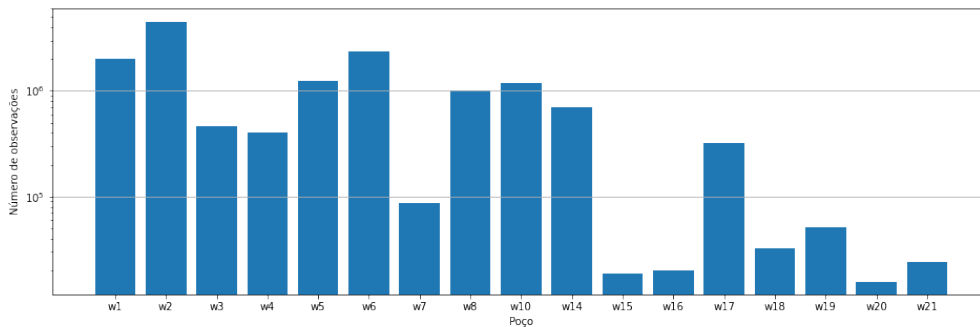


Figura 8 – Número de observações por conjunto de dados.

¹ Disponível em <https://data.mendeley.com/datasets/r7774rwc7v/1>

3.2 Análise da qualidade dos dados

Os conjuntos de dados apresentam lacunas e o primeiro tratamento executado foi a remoção de valores não numéricos e menores ou iguais a zero, pois esses valores não têm relevância física, uma vez que não são esperados essas condições em sistemas reais. Na figura 9 são apresentadas análises binárias para os 6 conjuntos de dados selecionados, nos quais as regiões claras indicam ausência de dados.

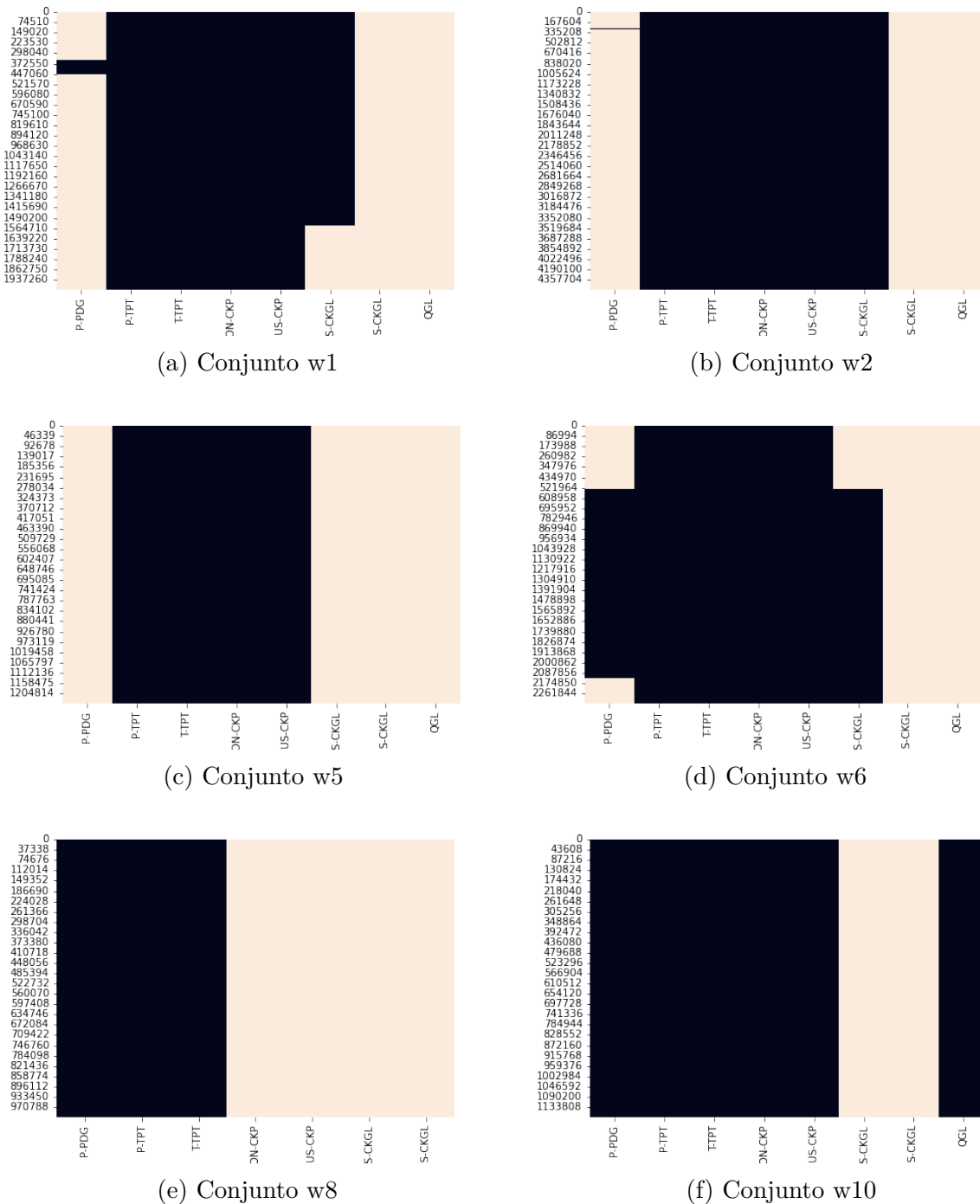


Figura 9 – Análise das lacunas de dados, nos quais regiões claras indicam ausência de dados. No eixo vertical são apresentados os índices das observações.

Pelas representações da figura 9 observa-se que o conjunto de dados w10 dispõe de dados das variáveis relevantes para esse trabalho e $1,1 \times 10^6$ observações, sendo esse o conjunto selecionado para ser utilizado.

Ao se analisar as distribuições dos valores das variáveis presentes no conjunto de dados w10, mostradas na figura 10, observa-se que as variáveis apresentam distribuições que não caracterizam a tendência para um valor específico e que assumem distribuições esperadas para os dados, indicando dados úteis para o escopo desse trabalho.

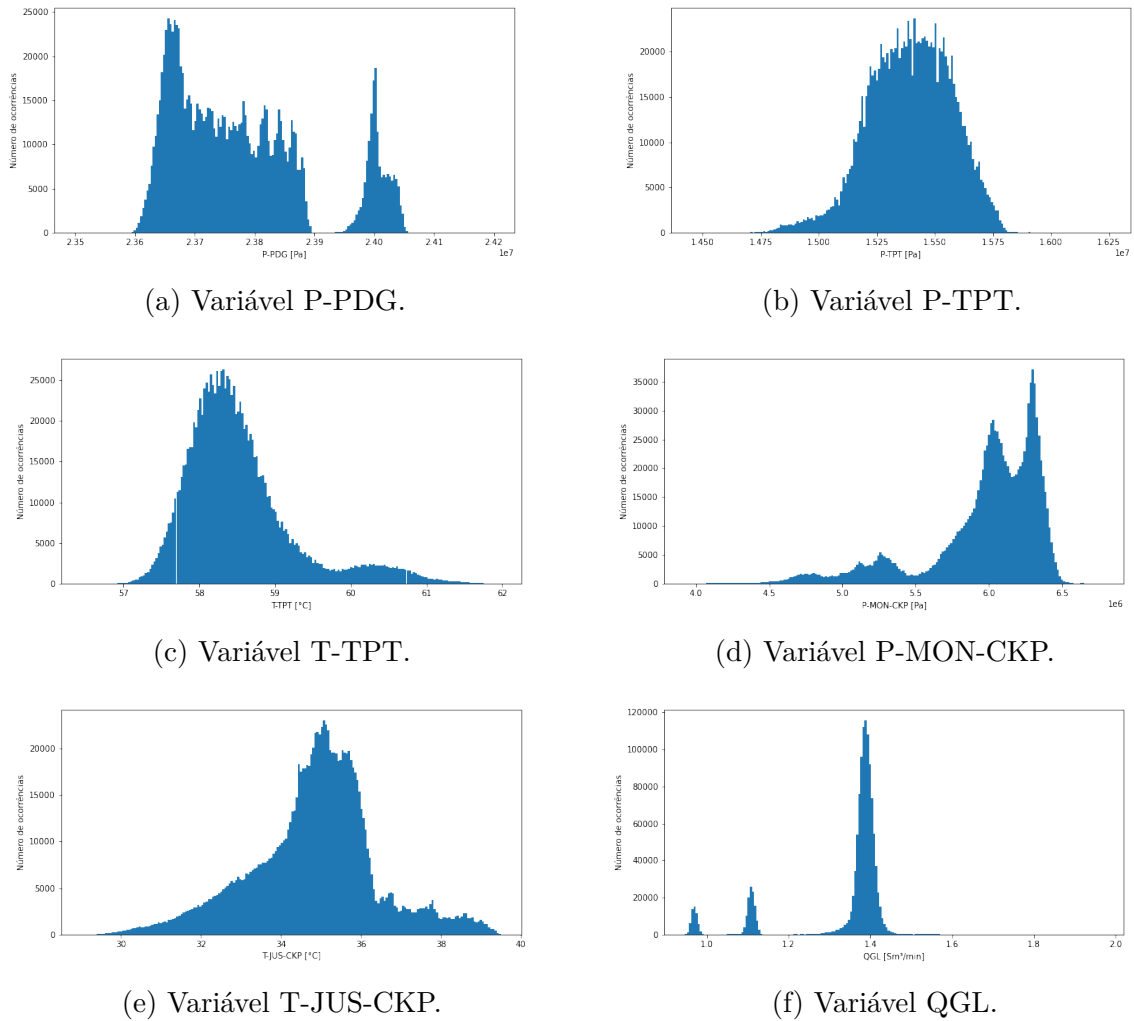


Figura 10 – Histograma das variáveis do conjunto de dados w10.

O conjunto de dados w6 também apresenta um número considerável de observações e poder-se-ia utilizar o intervalo intermediário que contempla as principais variáveis, porém esses dados P-PDG apresentam valores poucos dispersos, conforme é ilustrado na figura 11, indicando congelamento da variável P-PDG.

O comportamento das variáveis de w10 no período observado é apresentado na figura 12. O conjunto de dados é composto por dois períodos distintos (duas faixas pretas mostradas na figura 9-(f)), nos quais as variáveis têm comportamento nitidamente diferente,

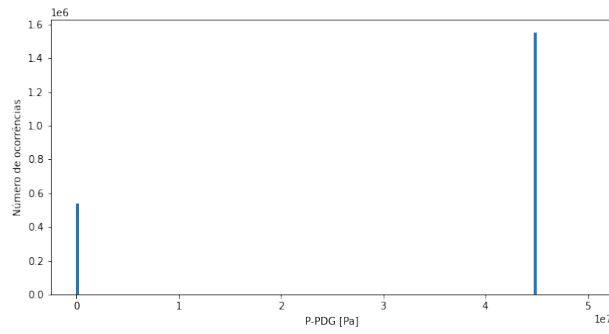


Figura 11 – Histograma da variável P-PDG do conjunto de dados w6.

possivelmente devido a alguma variação no processo. Os dados foram normalizados pela média no período para melhor apresentação.

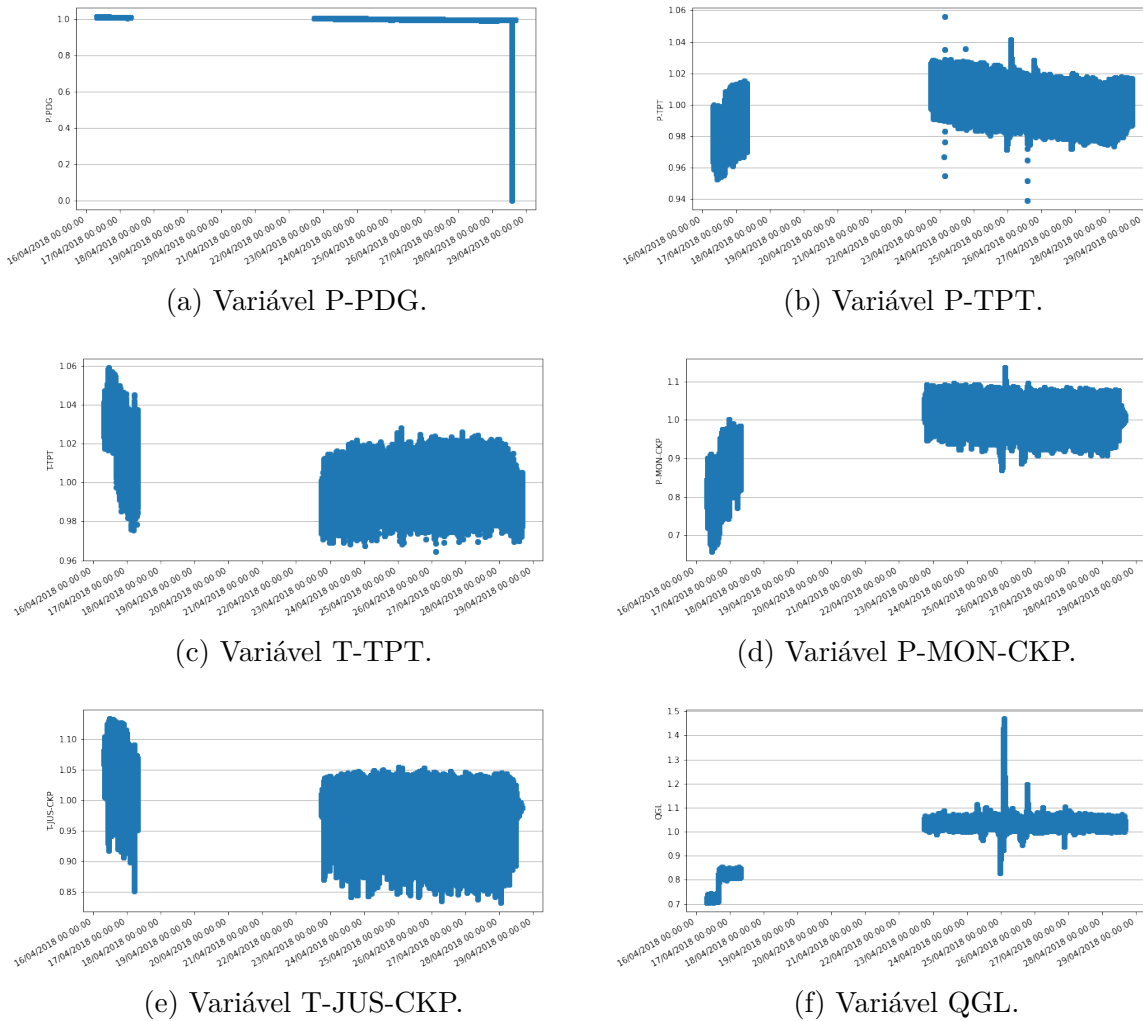


Figura 12 – Distribuição das variáveis do conjunto de dados w10 normalizadas pela média.

Pelo gráfico (b) da figura 12 pode-se observar dados espúrios que se destacam da tendência de dados próximos para a variável P-TPT. A variável P-PDG, gráfico (a) da figura 12, apresenta um comportamento anômalo próximo ao fim do período de observação.

Excluindo-se os pontos associados a esse comportamento da representação, pode-se notar de forma mais clara o comportamento da variável, conforme figura 13.

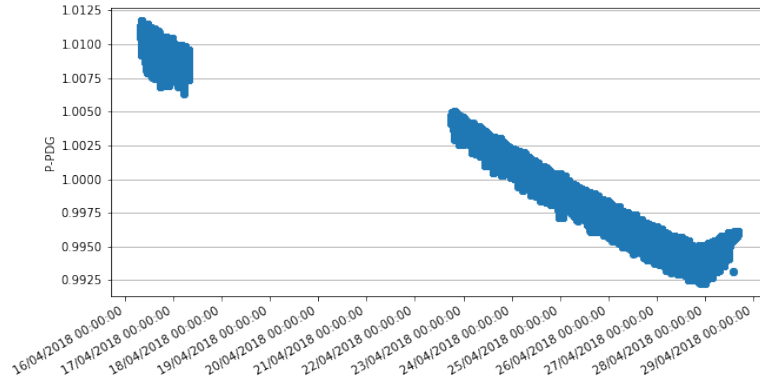


Figura 13 – Comportamento da variável P-PDG excluindo dados anômalos.

3.3 Análise da correlação entre variáveis

Uma vez que as variáveis estão tratando de um processo no qual o escoamento do fluido proveniente da rocha-reservatório se desloca desde o fundo do poço até a plataforma é possível que haja correlação entre as variáveis.

A análise multivariável trata de relações e, em particular, relações entre variáveis, tomadas em pares. A magnitude dessas relações é medida por coeficientes de correlação. A correlação entre duas variáveis é a covariância entre variáveis normalizadas (14). Para a normalização dos dados foi utilizada a técnica *Z-score*, que consiste na subtração de média de cada variável e divisão por seu desvio padrão.

Na figura 14 é apresentada a matriz de correlação entre as variáveis, com cores claras indicando alta correlação e cores escuras, baixa correlação.

Dessa matriz pode-se concluir que a variável T-TPT tem forte correlação inversa com P-MON-CKP e QGL e as variáveis P-MONT-CKP e QGL têm forte correlação direta. Essa análise leva à conclusão de que podemos reduzir o número de variáveis a serem consideradas e a figura 15 apresenta a matriz de correlação desconsiderando as variáveis T-TPT e P-MONT-CKP. Dessa forma as 4 variáveis restantes apresentam baixa correlação entre si.



Figura 14 – Matriz de correlação entre as 6 variáveis. Cores mais claras indicam alta correlação direta entre variáveis. Cores mais escuras indicam alta correlação inversa entre variáveis.

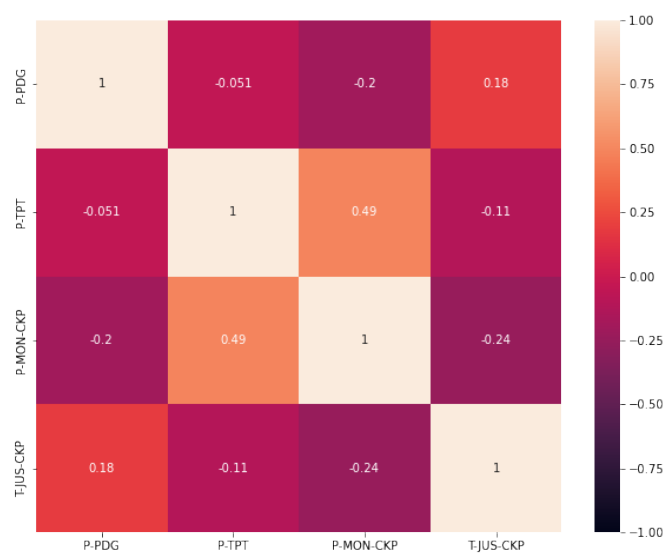


Figura 15 – Matriz de correlação entre as 4 variáveis. Cores mais claras indicam alta correlação direta entre variáveis. Cores mais escuras indicam alta correlação inversa entre variáveis.

3.4 Análise estatística dos dados

Para a análise dos principais dados estatísticos das variáveis é utilizada a ferramenta diagrama de caixa, ou *boxplot*, que apresenta os limites inferior e superior, quartis, mediana e dados discrepantes, ou *outliers*.

Conforme apresentado na figura 12 o conjunto de dados apresenta observações em dois momentos distintos e a análise será realizada para esses dois momentos.

Para todas as variáveis apresentadas na figura 16 estão presentes dados discrepantes e observa-se que há variação na mediana entre os dois momentos; porém, o intervalo entre o primeiro e o terceiro quartil são equivalentes, sendo esses dois períodos aqueles escolhidos para as análises.

3.5 Preparação dos dados

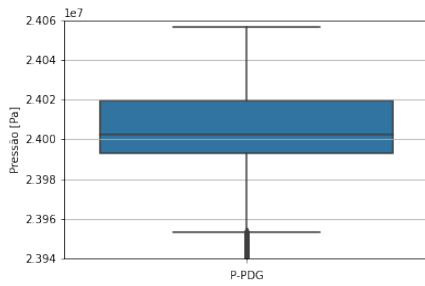
Os dados a serem utilizados na RNA precisam ser preparados. A primeira etapa é a remoção das variáveis que não são objetos das análises (timestamp, P-JUS-CKGL, T-JUS-CKGL, QGL, class, T-TPT).

Para permitir a análise dos resultados dos modelos de classificação (item 4.1) e de previsão (item 4.2) são separados três períodos de 20 minutos para compor os conjuntos de teste desses modelos. Os períodos foram selecionados para representar dados do começo, meio e fim do conjunto de dados. Esses conjuntos representam 0,34% dos dados totais.

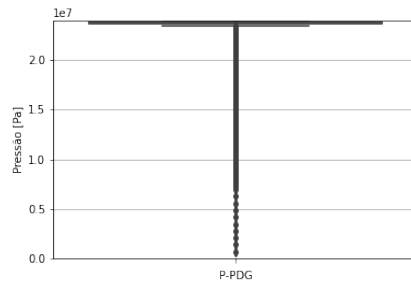
Os conjuntos de teste são separados dos demais dados de forma que os modelos em momento algum os observarão durante as etapas de treinamento e validação. Os gráficos das figuras 17, 18 e 19 apresentam o comportamento das quatro variáveis nos 3 períodos selecionados.

Os dados restantes são divididos em 75% para treinamento e 25% para validação dos modelos.

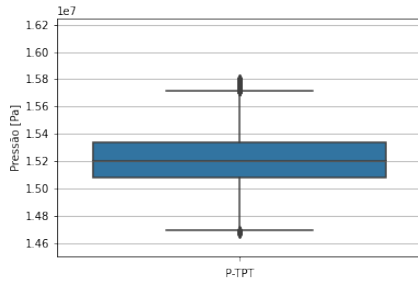
Cada uma das variáveis apresenta valores em ordem de grandeza consideravelmente diferentes e é aplicada a normalização *Z-score*, que consiste na subtração dos dados pela média e posterior divisão pelo desvio padrão. Os dados faltantes são substituídos pela média.



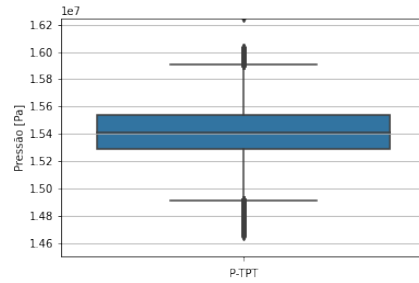
(a) P-PDG no primeiro momento.



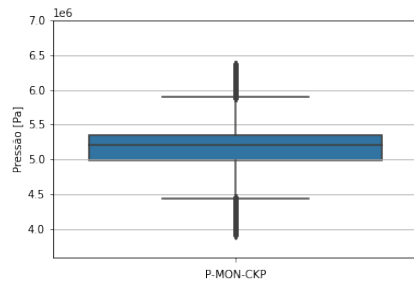
(b) P-PDG no segundo momento.



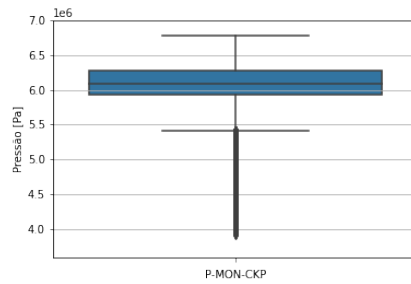
(c) P-TPT no primeiro momento.



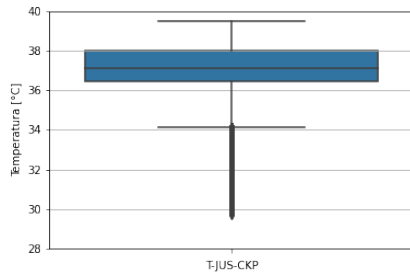
(d) P-TPT no segundo momento.



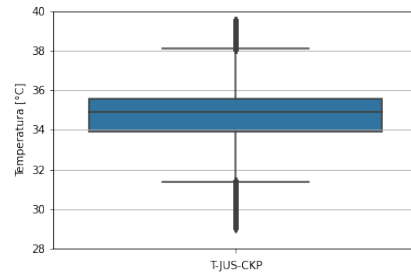
(e) P-MON-CKP no 1o. momento.



(f) P-MON-CKP no 2o. momento.



(g) T-JUS-CKP no 1o. momento.



(h) T-JUS-CKP no 2o. momento.

Figura 16 – Boxplot das variáveis.

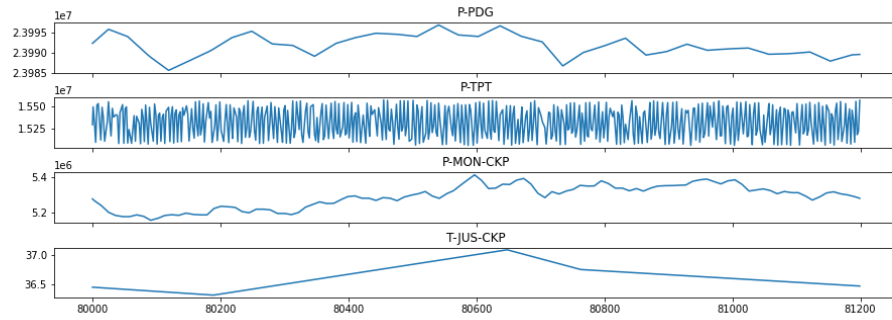


Figura 17 – Primeiro conjunto de teste.

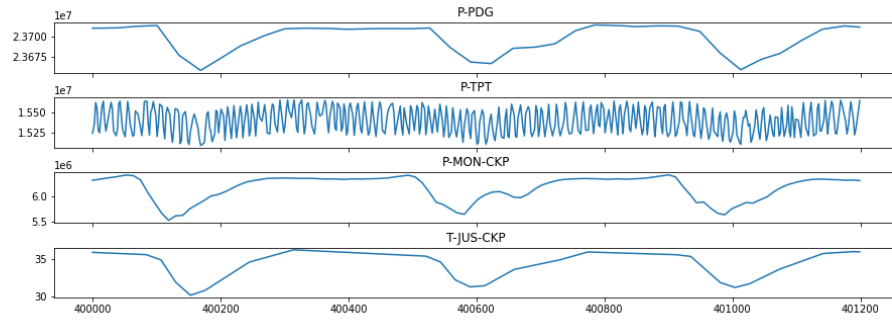


Figura 18 – Segundo conjunto de teste.

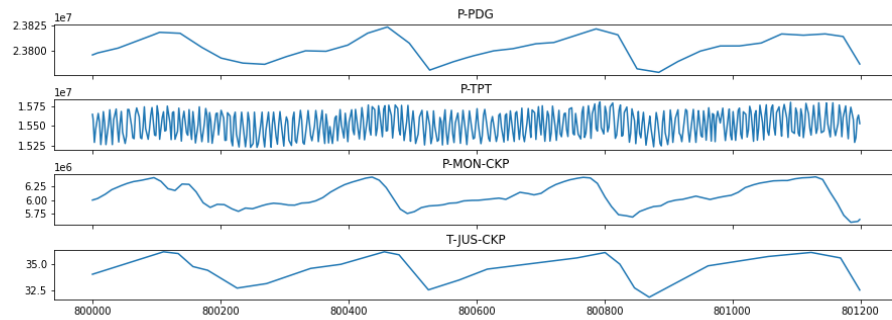


Figura 19 – Terceiro conjunto de teste.

4 DESENVOLVIMENTO E RESULTADOS

Na seção 4.1 é empregada uma rede neural artificial do tipo *Multilayer Perceptron* para classificar os dados da variável P-TPT e identificar anomalias. Os dados presentes no banco de dados 3W não possuem classificação quanto ao comportamento dos sensores, de forma que são geradas perturbações artificiais em parte dos dados da variável, associando rótulos a esses dados perturbados para treinamento e teste da rede neural artificial. No item 4.1 é detalhado o ruído gerado e suas características.

Uma forma alternativa para classificar dados de séries temporais é utilizar dados passados para estimar valores futuros e assim compará-los com os valores observados. Esse método é conhecido como *Sliding-window*. Na seção 4.2 é empregada essa técnica.

Os resultados obtidos com as duas técnicas são avaliados na seção 4.3 utilizando métricas específicas para cada método.

4.1 Modelo de Classificação

Na tarefa de classificação os dados são rotulados e é utilizado o aprendizado supervisionado com uma rede neural artificial do tipo *Multi Layer Perceptron*. A implementação dessa rede foi realizada através da biblioteca Keras do Tensorflow.

A rede implementada possui arquitetura equivalente à rede empregada no Saad et al (10), mantendo o mesmo número de camadas intermediárias e a proporcionalidade entre o número de neurônios da camada de entrada e das camadas intermediárias. Essa arquitetura é apresentada na figura 20, na qual as entradas $\mathbf{x} = (x_1, x_2, x_3, x_4)$ se referem às 4 variáveis a cada observação. Há três camadas escondidas, sendo a primeira com 8 neurônios, a segunda com 4 e a terceira com 2. Finalmente, a saída y indica a classe na qual a observação de entrada $\mathbf{x} = (x_1, x_2, x_3, x_4)$ se refere.

A camada de entrada e as camadas intermediárias utilizam a ativação ReLU 2.4. Na camada de saída é utilizada a função sigmoideal para retornar valores no intervalo $[0, 1]$, indicando valores próximos a 1 para dados da classe anômala e próximos a 0 para classe de dados normais, sendo o valor de 0,5 o limiar entre as duas classes.

Os dados disponibilizados no banco de dados não possuem classificação. Dessa forma, os dados da 3W são considerados como corretos (classe 0). São então incluídas perturbações em uma parte dos dados da variável P-TPT. Essas perturbações são geradas através da adição de um ruído gaussiano com média 0 e desvio padrão equivalente a $1,5\epsilon$, onde ϵ é a metade da amplitude máxima da variável previamente normalizada. Aplica-se esse ruído a 50% dos dados de treinamento e validação, de forma similar ao proposto por Guo e Runer (12), isto é, os dados com perturbação são considerados anômalos (classe 1).

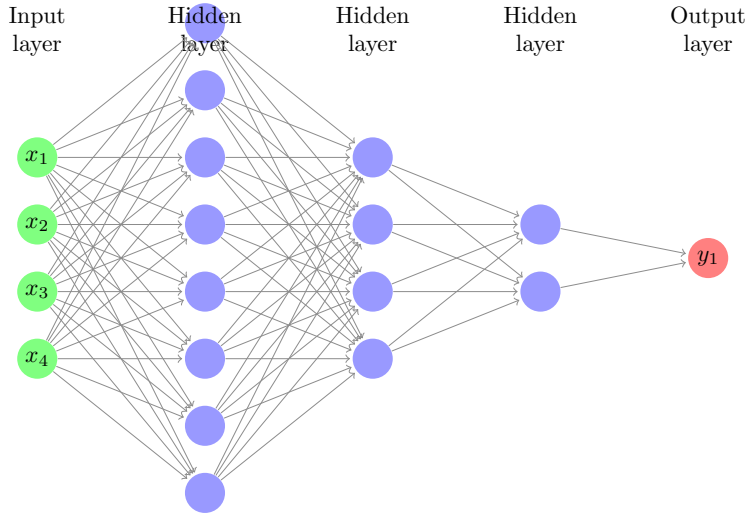


Figura 20 – Arquitetura de rede MLP utilizada.

Na figura 21 são apresentadas a variável P-TPT original e após a introdução da perturbação em um determinado período. Observa-se que em alguns pontos as duas curvas são coincidentes, ou seja, não foi aplicada perturbação. Essa perturbação introduzida caracteriza-se por um erro de deslocamento estocástico, conforme ilustrado na figura 2.

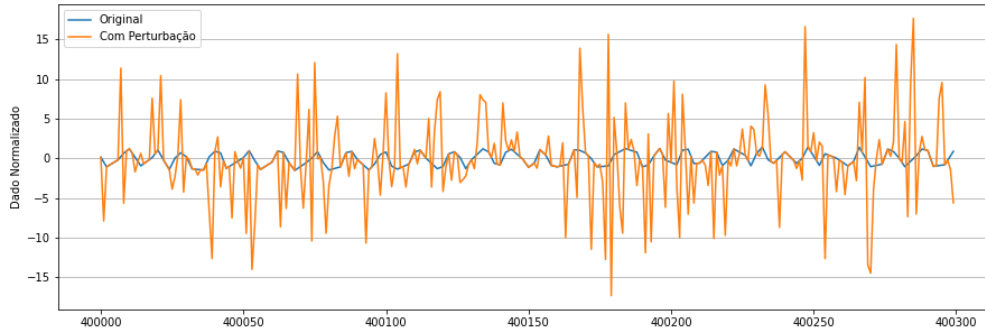


Figura 21 – Comportamento da variável P-TPT original (azul) e com perturbação introduzida (laranja).

O modelo MLP apresentado na figura 20 é configurado com taxa de aprendizado (otimizador Adam) e tamanho de *batch* conforme tabela 4. Cada um dos casos apresentados contempla pares de parâmetros (taxa de aprendizado e tamanho de *batch*) distintos para permitir uma análise comparativa e se identificar a melhor configuração para o problema em questão.

Em todos os casos o número de épocas no treinamento é limitado a 300, porém é utilizado um critério de parada em que o processo é interrompido caso a variação da perda em uma determinada época, em comparação com a vigésima época anterior, é igual ou inferior $1e^{-4}$. Esse atraso de 20 épocas, denominado *patience*, tem por objetivo evitar paradas em patamares intermediários. Na tabela 4 são apresentadas as épocas em que o

critério de parada foi atingido para cada caso.

A função perda utilizada é a entropia cruzada binária (*Binary Cross-Entropy*)

$$\text{Loss} = -\frac{1}{N} \sum_{i=1}^N y_i \log \hat{y}_i + (1 - y_i) \log(1 - \hat{y}_i), \quad (4.1)$$

onde $\hat{y}_i$ é o componente i da saída do modelo, y_i é o valor de referência e N é o número de saídas do modelo.

Os resultados são avaliados segundo as métricas Precisão, Revocação e *F1-score*, descritas na tabela 3, onde TP é o quantitativo de predições positivas verdadeiras, FP é o quantitativo de predições falso positivas e FN é o quantitativo de predições falso negativas. A métrica *F1-score* é incorporada pois combina a Precisão e Revocação de uma forma que permite a comparação entre diferentes configurações do modelo.

Tabela 3 – Métricas

Métrica	Definição	Cálculo
Precisão	Medida do percentual de predições positivas corretas	$\frac{TP}{TP+FP}$
Revocação	Medida do percentual de classificações positivas foram preditas corretamente	$\frac{TP}{TP+FN}$
<i>F1-score</i>	Média harmônica entre Precisão e Revocação	$\frac{2*Precisão*Revocação}{Precisão+Revocação}$

O resultado dessa análise é consolidado na tabela 4 e indica que as configurações dos casos 1 e 4 apresentam desempenho similar e superior aos demais, sendo o caso 1 considerado para as análises posteriores.

Tabela 4 – Resultados da Classificação

Caso	Taxa de aprendizado	Tamanho do batch	Épocas	Perda	Precisão	Revocação	F1-Score
1	$1e^{-3}$	120	194	0,239	0,988	0,843	0,910
2	$1e^{-4}$	120	138	0,242	0,981	0,845	0,908
3	$1e^{-5}$	120	300	0,254	0,979	0,830	0,898
4	$1e^{-3}$	60	190	0,237	0,983	0,846	0,910
5	$1e^{-3}$	180	198	0,239	0,981	0,847	0,909

4.2 Modelo de Previsão

Uma forma alternativa para se identificar anomalias é a comparação entre o dado obtido e a previsão de valor esperado.

Para prever valores em uma rede MLP um número fixo de observações prévias é considerado como entrada na rede em cada processo de treinamento, enquanto a saída são os valores preditos da série temporal, que é chamado de método *Sliding-Window* (15). A

figura 22 ilustra a aplicação do *Sliding-Window* em uma série temporal de uma variável, na qual os dados da janela temporal (*input window*) são as entradas de X_{t-1} a X_{t-p} da rede, onde p é o tamanho da janela. A saída da rede $\hat{X}_t$ é comparada com dados X_t para avaliação do erro.

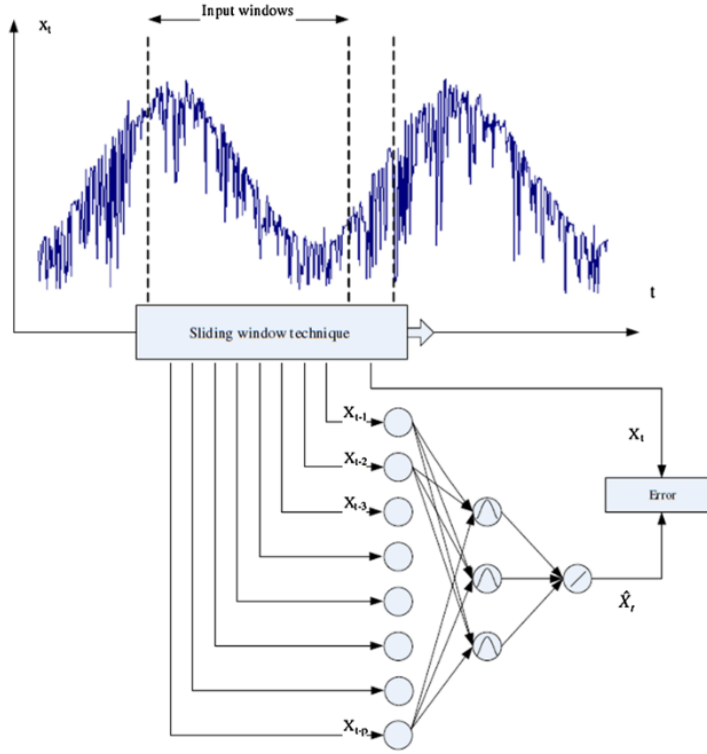


Figura 22 – Método *Sliding-Window* em uma RNA (15).

Nesse trabalho foi definido um tamanho de janela p de 3 minutos (180 observações) e a saída $\hat{X}_t$ é o valor de um instante da variável P-TPT. Como a rede tem um neurônio de saída, o método utiliza passo 1, ou seja, cada novo conjunto de dados de entrada se inicia com uma observação a frente na série temporal.

A configuração da rede MLP utilizada no modelo de classificação é empregada para a tarefa de previsão, com 720 neurônios na camada de entrada (4 variáveis e 180 observações), 1440, 720 e 360 neurônios na primeira, segunda e terceira camada intermediária, respectivamente, e 1 neurônio na camada de saída para representar a variável predita.

A camada de entrada e as camadas intermediárias utilizam a ativação ReLU (*Rectified Linear Unit*) e a camada de saída utiliza a ativação linear. É mantido o otimizador Adam com taxa de aprendizado inicial de $1e^{-7}$, conforme Saad et al (10).

A métrica utilizada para avaliar o modelo é o *Mean Absolute Percentage Error* (MAPE),

$$\text{MAPE}(y, \hat{y}) = \frac{1}{n_{\text{samples}}} \sum_{i=0}^{n_{\text{samples}}-1} \frac{|y_i - \hat{y}_i|}{\max(\epsilon, |y_i|)}, \quad (4.2)$$

aplicado no conjunto de validação, onde y_i representa a componente i o dado original, $\hat{y}_i$ a componente i do dado predito e ϵ é um número pequeno arbitrário (positivo) para evitar indeterminações quando y_i é nulo. O critério de parada de ajuste do modelo de delta mínimo para a perda de $1e^{-4}$, com *patience* de 20 épocas e limite de 1.000 épocas.

O resultado obtido com a taxa de aprendizado inicialmente proposta é apresentado na tabela 5, juntamente com uma análise de sensibilidade desse parâmetro, o que demonstra que a taxa de $1e^{-5}$ resultou em menor perda, sendo essa configuração empregada nos conjuntos de teste.

Tabela 5 – Resultados do modelo de previsão

Caso	Janela [s]	Taxa de aprendizado	Tamanho do batch	Épocas	Perda Validação
A	180	$1e^{-7}$	60	1000	78,7307
B	180	$1e^{-5}$	60	284	22,9253
C	180	$1e^{-4}$	60	145	35,1222

4.3 Avaliação dos Resultados

4.3.1 Modelo de Classificação

O modelo ajustado para tarefa de classificação é utilizado para classificar os conjuntos de teste previamente selecionados. Os conjuntos de teste são normalizados utilizando os mesmos parâmetros dos conjuntos de treinamento e validação. Como são apresentados dados originais ao modelo, espera-se que o resultado (classificação) apresente para todos os dados a classe 0, ou seja, que o modelo classifique esses dados como não anômalos. Uma vez que a função de saída da RNA é uma função sigmóide, o resultado é um valor contínuo entre 0 e 1, sendo considerado como 0,5 o limiar de separação entre classes 0 (normal) e 1 (anômala).

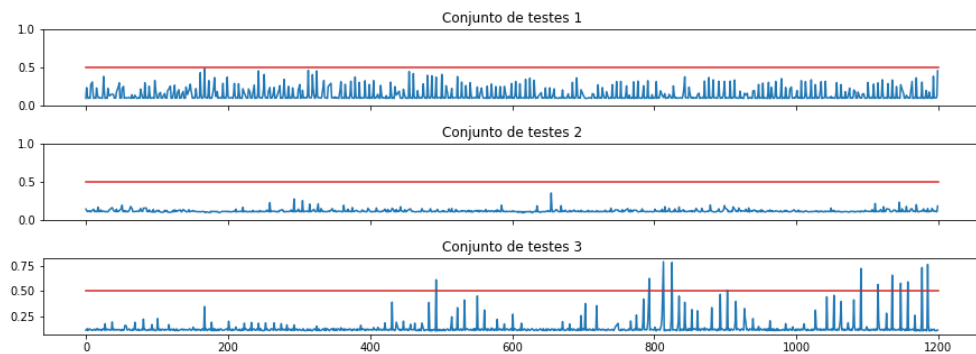


Figura 23 – Resultado da classificação dos conjuntos de testes. A linha azul representa o resultado do modelo de classificação e a linha vermelha, o limiar acima do qual se considera o dado como anômalo.

O modelo de classificação apresentou bom desempenho, resultando em erro de classificação somente para o conjunto de teste 3, para o qual 1,1% dos dados foram classificados erroneamente como anômalos. Esse bom desempenho pode ser oriundo da grande perturbação imposta à variável P-TPT, que gerou condições claramente distintas entre dados corretos e dados anômalos para a tarefa de treinamento. Um novo conjunto de dados foi gerado com a mesma forma de aplicar a perturbação, porém desvio padrão do ruído gaussiano foi reduzido pela metade. O efeito dessa redução pode ser observado na figura 24, que apresenta o dado original, a perturbação utilizada inicialmente e a perturbação reduzida.

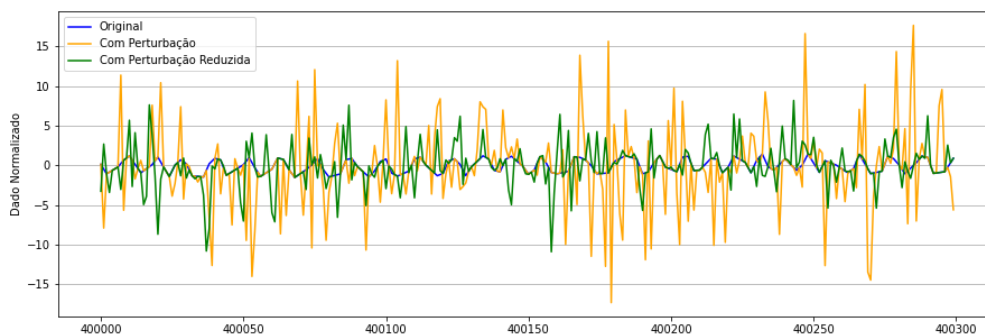


Figura 24 – Comportamento da variável P-TPT original (azul), variável P-TPT com perturbação inicial (laranja) e com perturbação reduzida (verde).

O modelo é novamente ajustado utilizando o conjunto de dados com perturbação reduzida e as configurações e resultados são apresentado na tabela 6.

Tabela 6 – Resultados da Classificação - Modelo com perturbação reduzida

Taxa de aprendizado	Tamanho do batch	Épocas	Perda	Precisão	Revocação	F1-Score
$1e^{-3}$	120	113	0,365	0,959	0,720	0,822

O desempenho do modelo é claramente afetado pela redução da perturbação, fazendo com que os dados corretos dos conjuntos de teste fossem em grande parte classificados como anômalos. O teste do modelo resultou em 100%, 5,5% e 3,9% de dados anômalos para os conjuntos de teste 1, 2 e 3, respectivamente.

Dessa análise conclui-se que o desempenho do modelo de classificação é dependente da discrepância entre dados corretos e dados anômalos, ou seja, pequenas variações no comportamento dos dados podem não ser detectados com o modelo empregado.

4.3.2 Modelo de Previsão

A análise do modelo de previsão ajustado consiste em utilizar os dados dos conjuntos de teste para prever o comportamento da variável em avaliação. Os conjuntos de teste

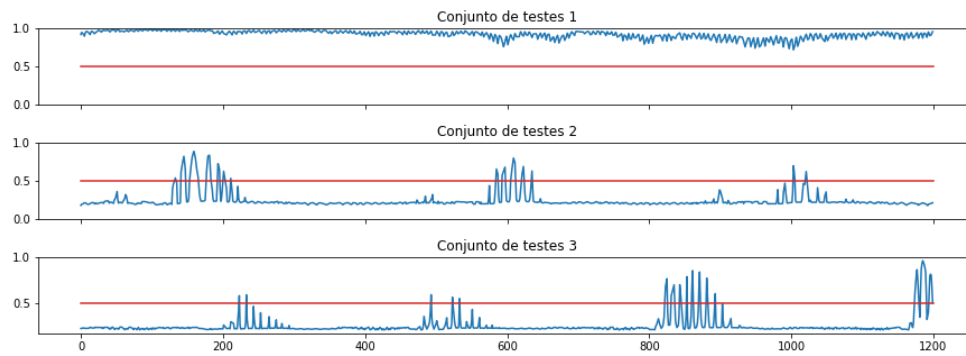


Figura 25 – Resultado da classificação dos conjuntos de testes. Linha azul representa a variável P-TPT normalizada, linha laranja o resultado do modelo de classificação e linha vermelha o limiar acima do qual se considera o dado como anômalo.

devem passar por processamento, conforme detalhado no item 4.2, para compor a entrada do modelo treinado.

Na figura 26 são apresentados os resultados do modelo para os três conjuntos de teste e indica aderência entre os dados originais e os dados preditos. Nos primeiros 180 pontos dos gráficos não há dados de previsão, pois essas observações foram utilizadas para prever os demais dados.

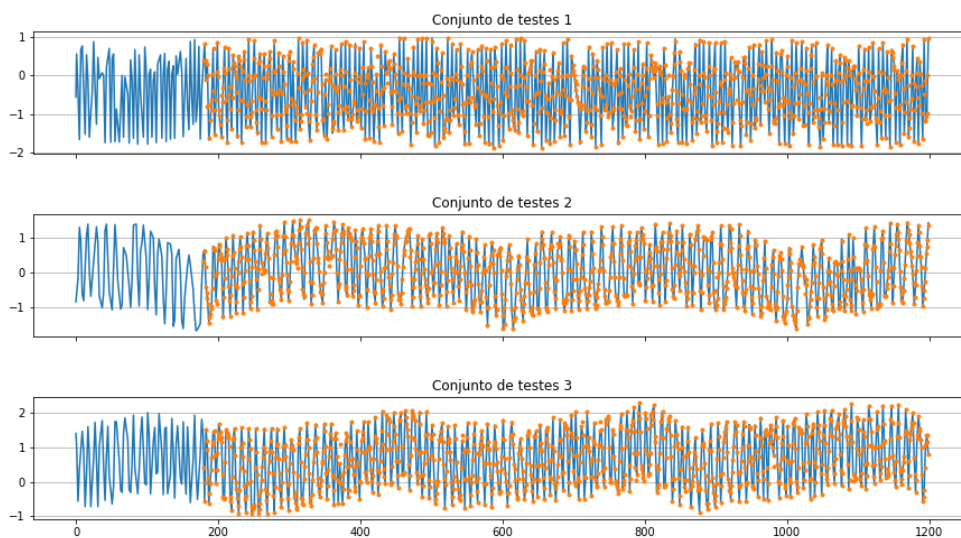


Figura 26 – Resultado do modelo de previsão aplicado aos conjuntos de testes 1, 2 e 3. Linha azul representa o dado original e os pontos em laranja o resultado do modelo.

De forma complementar aos gráficos apresentados são calculadas duas medidas de

erro para avaliação objetiva, coeficiente de determinação

$$R^2(y, \hat{y}) = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (4.3)$$

e raiz do erro quadrado médio

$$\text{RMSE}(y, \hat{y}) = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2}, \quad (4.4)$$

onde y_i representa o dado original, $\hat{y}_i$ o dado predito, $\bar{y} = \frac{1}{n} \sum_{i=1}^n y_i$ e n é o número de amostras. O coeficiente de determinação indica a qualidade do ajuste através da representação da proporção da variância de y pode ser explicada pelo modelo.

A tabela 7 apresenta os resultados do coeficiente de determinação para os três conjuntos de teste. Os resultados são similares e complementarmente foi avaliada a média dos erros quadrados. Essa métrica também indicou comportamento similar para os 3 conjuntos de teste.

Tabela 7 – Erro calculado para modelo de previsão

Métrica	Conjunto de Teste 1	Conjunto de Teste 2	Conjunto de Teste 3
R^2	0,951	0,979	0,971
$RMSE$	0,167	0,112	0,132

O modelo de previsão apresenta resultados de R^2 superior a 0,950, conseguindo reproduzir em grande parte os dados originais. Destaca-se que para o conjunto de testes 1, o modelo apresentou desempenho inferior quando comparado com os demais conjuntos de teste. Isso pode ser atribuído à menor variação da variável P-TPT com as demais variáveis nesse intervalo, fazendo com que o modelo tenha que predizer valores similares para condições distintas.

5 CONCLUSÃO

Esse trabalho investigou dois modelos de aprendizado de máquina supervisionado utilizando redes neurais artificiais para detectar dados anômalos de sensores de poços submarinos: o modelo de classificação e o modelo de previsão. Em ambos modelos foram utilizadas redes neurais artificiais do tipo *Multi-Layer Perceptron*.

O modelo de classificação utilizou dados com perturbação (ruído gaussiano) adicionada artificialmente para ajuste do modelo e foi testado com dados originais, classificando-os corretamente em mais de 98% dos casos. Ao se adicionar uma perturbação com desvio padrão equivalente a 50% do originalmente utilizado o modelo falhou em classificar 100% dos dados de um conjunto de teste. Com isso, identificou-se que o desempenho do modelo proposto é dependente da perturbação adicionada ao conjunto de treinamento e que o uso de ruído artificial é um fator preponderante para esse desempenho.

No modelo de previsão foi empregada a técnica *Sliding-Window*, sendo utilizado um trecho (janela) da série temporal para prever dados subsequentes. A janela utilizada foi de 180 segundos, prevendo o dado do segundo seguinte. O modelo proposto teve êxito em prever os dados dos conjuntos de teste e os coeficientes de determinação R^2 calculados foram superiores a 0,95.

Dentre os dois modelos propostos, o modelo de previsão tem maior potencial de aplicação prática, pois permite que seja estabelecido um critério de desvio aceitável para a variável em análise e assim seja caracterizado um dado medido como anômalo caso ultrapasse esse desvio a partir do dado predito. Contudo, esse modelo tem um custo computacional comparativamente maior, com tempo de processamento cerca de 5 vezes superior.

O banco de dados analisados possui 10^6 observações das variáveis, o que representa aproximadamente 12 dias de observações registradas com frequência de 1 Hz. Esses dados foram suficientes para treinamento, validação e teste dos modelos propostos.

As variáveis analisadas possuem uma componente de baixa frequência associada ao processo ao qual estão inseridas e uma componente de mais alta frequência possivelmente associada a ruído. Para trabalhos futuros recomenda-se excluir essas componentes de alta frequência ou aplicar algum filtro nos dados de forma a ajustar o modelo de previsão a dados menos ruidosos com o objetivo de melhorar o seu desempenho. Com relação ao modelo de classificação, se faz necessário o uso de dados reais rotulados para melhor avaliação e um novo trabalho comparativo entre modelos de classificação e previsão pode ser desenvolvido.

REFERÊNCIAS

- 1 VARGAS, R. E. V. *et al.* A realistic and public dataset with rare undesirable real events in oil wells. **Journal of Petroleum Science and Engineering**, v. 181, p. 106223, out. 2019. ISSN 09204105. Disponível em: <https://linkinghub.elsevier.com/retrieve/pii/S0920410519306357>.
- 2 SOBRINHO, E. d. S. P. *et al.* Uma ferramenta para detectar anomalias de produção utilizando aprendizagem profunda e árvore de decisão. *In: .* Rio de Janeiro: IBP, 2020. Disponível em: https://icongresso.ibp.itarget.com.br/arquivos/trabalhos_completos/ibp/3/final.IBP0938_20_15032020_222247.pdf.
- 3 ANJOS, J. L. *et al.* Digital Twin for Well Integrity with Real Time Surveillance. *In: Day 2 Tue, May 05, 2020.* Houston, Texas, USA: OTC, 2020. p. D021S018R001. Disponível em: <https://onepetro.org/OTCONF/proceedings/20OTC/2-20OTC/Houston,%20Texas,%20USA/110954>.
- 4 JAGER, G. *et al.* Assessing neural networks for sensor fault detection. *In: 2014 IEEE International Conference on Computational Intelligence and Virtual Environments for Measurement Systems and Applications (CIVEMSA).* Ottawa, ON, Canada: IEEE, 2014. p. 70–75. ISBN 9781479926145 9781479926138. Disponível em: <https://ieeexplore.ieee.org/document/6841441>.
- 5 TEH, H. Y.; KEMPA-LIEHR, A. W.; WANG, K. I.-K. Sensor data quality: a systematic review. **Journal of Big Data**, v. 7, n. 1, p. 11, dez. 2020. ISSN 2196-1115. Disponível em: <https://journalofbigdata.springeropen.com/articles/10.1186/s40537-020-0285-1>.
- 6 YU, Y. *et al.* Time Series Outlier Detection Based on Sliding Window Prediction. **Mathematical Problems in Engineering**, v. 2014, p. 1–14, 2014. ISSN 1024-123X, 1563-5147. Disponível em: <http://www.hindawi.com/journals/mpe/2014/879736/>.
- 7 VACHTSEVANOS, G. J. (ed.). **Intelligent fault diagnosis and prognosis for engineering systems.** Hoboken, N.J: Wiley, 2006. OCLC: ocm64442758. ISBN 9780471729990.
- 8 HODGE, V.; AUSTIN, J. A Survey of Outlier Detection Methodologies. **Artificial Intelligence Review**, v. 22, n. 2, p. 85–126, out. 2004. ISSN 0269-2821. Disponível em: <http://link.springer.com/10.1023/B:AIRE.0000045502.10941.a9>.
- 9 GUPTA, M. *et al.* Outlier Detection for Temporal Data: A Survey. **IEEE Transactions on Knowledge and Data Engineering**, v. 26, n. 9, p. 2250–2267, set. 2014. ISSN 1041-4347. Disponível em: <http://ieeexplore.ieee.org/document/6684530/>.
- 10 SAAD, A. M. *et al.* Using Neural Network Approaches to Detect Mooring Line Failure. **IEEE Access**, v. 9, p. 27678–27695, 2021. ISSN 2169-3536. Disponível em: <https://ieeexplore.ieee.org/document/9352003/>.
- 11 HAYKIN, S. S. **Neural networks and learning machines.** 3rd ed. ed. New York: Prentice Hall, 2009. OCLC: ocn237325326. ISBN 978-0-13-147139-9.

- 12 GUO, T.-H.; NURRE, J. Sensor failure detection and recovery by neural networks. *In: IJCNN-91-Seattle International Joint Conference on Neural Networks*. Seattle, WA, USA: IEEE, 1991. i, p. 221–226. ISBN 9780780301641. Disponível em: <http://ieeexplore.ieee.org/document/155180/>.
- 13 DAVIES, D. R.; AGGREY, G. H. Tracking the state and diagnosing Down Hole Permanent Sensors in Intelligent Well Completions with Artificial Neural Network. *In: All Days*. Aberdeen, Scotland, U.K.: SPE, 2007. p. SPE-107198-MS. Disponível em: <https://onepetro.org/SPEOE/proceedings/07OE/All-07OE/SPE-107198-MS/141966>.
- 14 BARTHOLOMEW, D. Analysis and Interpretation of Multivariate Data. *In: International Encyclopedia of Education*. Elsevier, 2010. p. 12–17. ISBN 9780080448947. Disponível em: <https://linkinghub.elsevier.com/retrieve/pii/B9780080448947013038>.
- 15 VAFAEIPOUR, M. *et al.* Application of sliding window technique for prediction of wind velocity time series. **International Journal of Energy and Environmental Engineering**, v. 5, n. 2-3, p. 105, jul. 2014. ISSN 2008-9163, 2251-6832. Disponível em: <http://link.springer.com/10.1007/s40095-014-0105-5>.