

MESTRADO

CIÊNCIAS EMPRESARIAS

TRABALHO FINAL DE MESTRADO

DISSERTAÇÃO

**PREVISÃO DOS PREÇOS DE LICENÇAS DE EMISSÕES DE
CARBONO, RECORRENDO A MÉTODOS DE INTELIGÊNCIA
ARTIFICIAL**

JOÃO DIOGO NUNES GUERREIRO DOS SANTOS

Junho 2025

MESTRADO EM CIÊNCIAS EMPRESARIAIS

TRABALHO FINAL DE MESTRADO **DISSERTAÇÃO**

**PREVISÃO DOS PREÇOS DE LICENÇAS DE EMISSÕES DE
CARBONO, RECORRENDO A MÉTODOS DE INTELIGÊNCIA
ARTIFICIAL**

JOÃO DIOGO NUNES GUERREIRO DOS SANTOS

Orientador:

**PROF. DOUTOR JOSÉ MIGUEL ARAGÃO CELESTINO SOARES
PROF. DOUTOR RICARDO SIMÕES SANTOS**

Junho 2025

“Imagination is more important than knowledge. For knowledge is limited, whereas imagination embraces the entire world, stimulating progress, giving birth to evolution.”

(Albert Einstein, 1929)

AGRADECIMENTOS

A realização desta tese representou o culminar de um percurso académico longo e exigente, marcado por desafios, mudanças e algumas pausas inevitáveis. Este trabalho só foi possível graças ao apoio, incentivo e generosidade de várias pessoas a quem gostaria de expressar a minha mais profunda gratidão.

Em primeiro lugar, ao Professor Doutor Ricardo Simões Santos, pela disponibilidade, orientação e valiosas sugestões ao longo de todo o processo. A sua contribuição, em particular no que diz respeito às fontes de dados utilizadas, foi fundamental para o desenvolvimento deste trabalho.

Ao Professor Doutor José Miguel Aragão Celestino Soares, pela constante disponibilidade e por nunca deixar de me recordar — com firmeza e gentileza — da importância de concluir este projeto. A sua presença e incentivo foram sempre uma motivação.

À minha namorada, que sempre acreditou em mim, mesmo nos momentos em que eu próprio duvidava. A ela, que me apoiou incondicionalmente e a quem tive de roubar algum tempo e atenção para me dedicar a esta etapa, deixo um agradecimento cheio de carinho.

À minha família, pelo apoio inabalável e pela oportunidade que me deram de seguir este percurso académico. À minha mãe, que sempre defendeu que estudar é um dos caminhos mais seguros para crescer e que nunca deixou de me apoiar. E ao meu pai, que já não está entre nós, mas cuja fé em mim e no meu potencial permanece viva e presente em tudo o que faço.

Aos meus colegas de trabalho, que lidaram com algumas oscilações de humor ao longo desta jornada — nem sempre foi fácil — mas que, mesmo assim, estiveram sempre disponíveis para ajudar e facilitar a minha vida. O vosso apoio foi, sem dúvida, essencial.

Por fim, a mim próprio. A mim, que concluí a parte letiva deste mestrado em 2020, mas que, por entre imprevistos, mudanças e episódios menos felizes, fui adiando a conclusão desta etapa. A mim, que agora, finalmente, a termino. Porque nunca é tarde para concluir aquilo que um dia começámos.

A todos, o meu sincero obrigado.

GLOSSÁRIO

ADF – *Augmented Dickey-Fuller*

AGC – Arquitetura de Geração Convolutacional

CO₂ – Dióxido de Carbono

CNN – Rede Neural Convolutacional

EU ETS – Sistema de Comércio de Licenças de Emissão da União Europeia

GEE – Gases com Efeito de Estufa

IA – Inteligência Artificial

LSTM – Memória de Curto e Longo Prazo

MAE – Erro Absoluto Médio

ML – Aprendizagem Automática

MSR – Reserva de Estabilidade do Mercado

R² – Coeficiente de Determinação

RMSE – Raiz do Erro Quadrático Médio

RNA – Rede Neural Artificial

RNR – Rede Neural Recorrente

RESUMO

Esta dissertação avalia a capacidade preditiva dos preços do carbono no âmbito do Sistema de Comércio de Licenças de Emissão da União Europeia (EU ETS), recorrendo a técnicas de machine learning. A análise foca-se na previsão do nível absoluto de preços diários, em vez dos retornos, permitindo uma avaliação mais direta da aplicabilidade dos modelos a contextos reais de decisão. Foram testadas diversas abordagens, com destaque para modelos *Random Forest* aplicados a diferentes regimes do mercado (Fase 3, Fase 4), bem como uma estratégia de janela deslizante, que permite adaptar o modelo às dinâmicas mais recentes do mercado.

Os resultados demonstram que modelos simples, quando corretamente ajustados, podem superar arquiteturas complexas como a CNN-LSTM, especialmente em contextos com dados tabulares e regimes instáveis. A abordagem de janela deslizante revelou-se particularmente eficaz, reforçando a importância da adaptabilidade em mercados caracterizados por mudanças estruturais frequentes. Conclui-se que, apesar dos desafios inerentes à previsão diária, é possível melhorar significativamente a performance preditiva com técnicas adequadas e segmentação temporal, contribuindo para uma melhor compreensão do comportamento dos preços no EU ETS.

Palavras-chave: Sustentabilidade, Sistema de Comércio de Licenças de Emissão da União Europeia (EU ETS), Preço do carbono, Previsão de preços, *Machine Learning*, *Random Forest*, *Deep learning*, CNN-LSTM

ABSTRACT

This dissertation assesses the predictive power of carbon price forecasting within the European Union Emissions Trading System (EU ETS), using machine learning techniques. Unlike prior studies focusing on returns, this work forecasts the daily absolute price levels, enabling a more practical evaluation of model performance. Several approaches were tested, including Random Forest models trained on different phases of the market (Phase 3, Phase 4), and a sliding window strategy designed to adapt to recent market dynamics.

The results show that simpler models, when properly tuned, can outperform complex architectures such as CNN-LSTM, particularly when dealing with tabular data and regime-dependent patterns. The sliding window method proved especially effective, highlighting the importance of adaptability in markets characterised by frequent structural changes. Despite the inherent challenges of daily forecasting, the study demonstrates that predictive accuracy, can be substantially improved through careful model design and temporal segmentation, offering a more realistic perspective on carbon price behaviour in the EU ETS.

Keywords: Sustainability, Carbon prices, European Union Emissions Trading System (EU ETS), Machine learning, Random Forest, Price forecasting, Adaptive models, Deep learning, CNN-LSTM

AGRADECIMENTOS	I
GLOSSÁRIO	II
RESUMO	III
LISTA DE TABELAS.....	VI
LISTA DE FIGURAS	VI
LISTA DE ANEXOS.....	VI
1. INTRODUÇÃO.....	1
2. REVISÃO DE LITERATURA	2
2.1 SUSTENTABILIDADE	2
2.1.1 <i>Origem, conceitos e importância.....</i>	2
2.1.2 <i>O papel das emissões de CO₂ na garantia da sustentabilidade.....</i>	4
2.2 MERCADOS DE LICENÇAS DE EMISSÕES DE CO ₂	5
2.2.1 <i>Propósito e Funcionamento.....</i>	5
2.2.2 <i>Principais Mercados Existentes</i>	7
2.2.3 <i>Desafio Existentes.....</i>	9
2.3 MACHINE LEARNING.....	10
2.3.1 <i>O que é Machine Learning</i>	10
2.3.2 <i>Funções Fundamentais em Machine Learning.....</i>	11
2.3.3 <i>Diferenças entre Machine Learning tradicional e Deep Learning.....</i>	12
2.3.4 <i>Como os modelos aprendem e funcionam</i>	13
2.3.5 <i>Tipos de problemas e áreas de aplicabilidade</i>	15
2.3.6 <i>Diferenças face aos modelos econométricos tradicionais.....</i>	16
2.4 MODELOS UTILIZADOS NA PREVISÃO DE PREÇOS	16
2.4.1 <i>Features utilizadas na definição dos modelos.....</i>	16
2.4.2 <i>Modelos Tradicionais de Previsão.....</i>	18
2.4.3 <i>Técnicas de Machine Learning e Deep Learning.....</i>	19
3. METODOLOGIA E DEFINIÇÃO DOS DADOS.....	20
3.1 CARACTERIZAÇÃO DO ESTUDO	20
3.1.2 <i>ESTRATÉGIA DE INVESTIGAÇÃO ADOTADA</i>	21
3.1.3 <i>MÉTODO DE RECOLHA DE DADOS</i>	21
3.2 INTRODUÇÃO AO CASO DE ESTUDO	22
4. DEFINIÇÃO DO MODELO.....	22
4.1 ESCOLHA DAS <i>FEATURES</i>	23
4.2 PRÉ-PROCESSAMENTO DE DADOS E TRANSFORMAÇÕES	24
4.2.1 <i>Interpolação dos Dados.....</i>	24
4.2.2 <i>Transformações</i>	25
4.2.3 <i>Outliers e Correlações.....</i>	26
4.2.4 <i>Normalização dos Dados.....</i>	26
4.3 TÉCNICAS DE MACHINE LEARNING E DEEP LEARNING ADOTADAS.....	27
4.4 MÉTRICAS DE DESEMPENHO.....	28
4.5 ESTRATÉGIAS DE ANÁLISE E VALIDAÇÃO	29
4.5.1 <i>Análise Inter-Regime e Intra-Regime</i>	29
4.5.2 <i>Seleção de Features em Duas Etapas.....</i>	30
4.5.3 <i>Estratégia de Janela Deslizante (Moving Window).....</i>	30
5. APRESENTAÇÃO E ANÁLISE DOS RESULTADOS	31
6. CONCLUSÕES, LIMITAÇÕES E RECOMENDAÇÕES DE TRABALHO FUTURO	34
7. REFERÊNCIAS BIBLIOGRÁFICAS	36
8. ANEXOS.....	40

LISTA DE TABELAS

TABELA I - TABELA COM OS RESULTADOS DAS MÉTRICAS RMSE, MAE E R2 PARA OS MODELOS TESTADOS	33
TABELA II - TABELA COM O RESUMO DAS FUNÇÕES FUNDAMENTAIS DO MACHINE LEARNING.....	40
TABELA III - TABELA QUE RESUME AS PRINCIPAIS DIFERENÇAS ENTRE O ML TRADICIONAL E O DL.....	40
TABELA IV - TABELA COM OS TIPOS DE APRENDIZAGEM E AS CATEGORIAS DE PROBLEMAS EM ML.....	40
TABELA V - TABELA QUE RESUME AS DIFERENÇAS ENTRE MODELOS ECONOMETRICOS E MODELOS DE ML.....	40
TABELA VI - TABELA COM A DESCRIÇÃO DE CADA VARIÁVEL E A QUE CATEGORIA PERTENCEM	41
TABELA VII - TABELA COM A ASSIMETRIA POR VARIÁVEL E TRANSFORMAÇÃO APLICADA EM FUNÇÃO DO RESULTADO DO TESTE	41
TABELA VIII - TABELA COM O SUMÁRIO ESTATÍSTICO DAS VARIÁVEIS QUE IRÃO SERVIR COMO BASE DO MODELO PREDITIVO	42

LISTA DE FIGURAS

FIGURA 1 - DISTRIBUIÇÃO BRENT_PRICE ANTES VS DEPOIS	42
FIGURA 2- DISTRIBUIÇÃO CARBON_PRICE ANTES VS DEPOIS.....	42
FIGURA 3 - DISTRIBUIÇÃO CO2_MT ANTES VS DEPOIS	43
FIGURA 4 - DISTRIBUIÇÃO COAL_PRICE ANTES VS DEPOIS.....	43
FIGURA 5 - DISTRIBUIÇÃO ECB_RATE ANTES VS DEPOIS.....	43
FIGURA 6 - DISTRIBUIÇÃO EDPR_PRICE ANTES VS DEPOIS.....	43
FIGURA 7 - DISTRIBUIÇÃO EU_AVG_SPOT_PRICE ANTES VS DEPOIS	44
FIGURA 8 - DISTRIBUIÇÃO EUROSTOXX_PRICE ANTES VS DEPOIS.....	44
FIGURA 9 - DISTRIBUIÇÃO FOSSIL_MW ANTES VS DEPOIS.....	44
FIGURA 10 - DISTRIBUIÇÃO GAS_PRICE ANTES VS DEPOIS	44
FIGURA 11 - DISTRIBUIÇÃO GERMAN_2Y_BOND ANTES VS DEPOIS.....	45
FIGURA 12 - DISTRIBUIÇÃO GERMAN_5Y_BOND ANTES VS DEPOIS.....	45
FIGURA 13 - DISTRIBUIÇÃO GERMAN_30Y_BOND ANTES VS DEPOIS.....	45
FIGURA 14 - DISTRIBUIÇÃO GPRD ANTES VS DEPOIS.....	45
FIGURA 15 - DISTRIBUIÇÃO INFLATION ANTES VS DEPOIS	46
FIGURA 16 - DISTRIBUIÇÃO NUCLEAR_ENERGY_INDEX ANTES VS DEPOIS	46
FIGURA 17 - DISTRIBUIÇÃO NUCLEAR_MW ANTES VS DEPOIS	46
FIGURA 18 - DISTRIBUIÇÃO RENEWABLE_MW ANTES VS DEPOIS	46
FIGURA 19 - DISTRIBUIÇÃO SOLAR_ENERGY_INDEX ANTES VS DEPOIS	47
FIGURA 20 - DISTRIBUIÇÃO VIX_PRICE ANTES VS DEPOIS.....	47
FIGURA 21- ESTRUTURA DA CNN-LSTM USADA PARA TESTAR A PREVISÃO DO PREÇO DO EU ETS	47
FIGURA 22 - ESTRUTURA DO RANDOM FOREST REGRESSOR USADO PARA TESTAR A PREVISÃO DO PREÇO DO EU ETS	48
FIGURA 23 - ESTRUTURA DO RANDOM FOREST REGRESSOR COM JANELA DESLIZANTE PARA A PREVISÃO DO EU ETS	48

LISTA DE ANEXOS

ANEXO I – FUNÇÕES FUNDAMENTAIS EM MACHINE LEARNING	40
ANEXO II – MACHINE LEARNING VS DEEP LEARNING	40
ANEXO III – TIPOS DE APRENDIZAGEM E CLASSIFICAÇÃO DOS PROBLEMAS EM MACHINE LEARNING	40
ANEXO IV – DIFERENÇAS FUNDAMENTAIS ENTRE OS MODELOS TRADICIONAIS ECONOMETRICOS E O MACHINE LEARNING	40
ANEXO V – VARIÁVEIS ESCOLHIDAS COM BASE NA LITERATURA, PARA A CONSTRUÇÃO DO MODELO DE PREVISÃO	41
ANEXO VI – DIAGNOSTICO DE ASSIMETRIA E TRANSFORMAÇÕES APLICADAS POR VARIÁVEL	41

ANEXO VII – SUMÁRIO ESTATÍSTICO DOS DADOS JÁ TRATADOS E PRONTO A UTILIZAR NO MODELO. DISTRIBUIÇÕES DE CADA VARIÁVEL ANTES E DEPOIS DAS TRANSFORMAÇÕES.	41
ANEXO VIII – CÓDIGO PYTHON DOS MODELOS USADOS.....	47
ANEXO IX – FUNÇÃO <i>LOSS</i> AO LONGO DAS VÁRIAS <i>EPOCHS</i> PARA O TREINO, VALIDAÇÃO E TESTE NO MATLAB USANDO UMA REDE NEURONAL SIMPLES COM 100 CAMADAS.....	48
ANEXO X – FUNÇÃO <i>LOSS</i> AO LONGO DAS VÁRIAS <i>EPOCHS</i> PARA O TREINO, VALIDAÇÃO E TESTE NO <i>JUPYTER NOTEBOOK</i> , BIBLIOTECA <i>TENSORFLOW</i> DO <i>PYTHON</i> USANDO UMA REDE NEURONAL DO TIPO CNN-LSTM SEQ2SEQ	49
ANEXO XI – GRÁFICO QUE MOSTRA O RESULTADO DO PREÇO REAL VS A PREVISÃO DO PREÇO USANDO A CNN-LSTM	49
ANEXO XII – GRÁFICO QUE MOSTRA O RESULTADO DO PREÇO RAL VS A PREVISÃO DO PREÇO USANDO O <i>RANDOM FOREST REGRESSOR</i> COM AMOSTRA COMPLETA	49
ANEXO XIII – GRÁFICO QUE MOSTRA O RESULTADO DO PREÇO RAL VS A PREVISÃO DO PREÇO USANDO O <i>RANDOM FOREST REGRESSOR</i> COM DADOS DA FASE 3.....	50
ANEXO XIV – GRÁFICO QUE MOSTRA O RESULTADO DO PREÇO RAL VS A PREVISÃO DO PREÇO USANDO O <i>RANDOM FOREST REGRESSOR</i> COM DADOS DA FASE 4.....	50
ANEXO XV – GRÁFICO QUE MOSTRA O RESULTADO DO PREÇO RAL VS A PREVISÃO DO PREÇO USANDO O A JANELA DESLIZANTE NO <i>RANDOM FOREST REGRESSOR</i>	50

1. INTRODUÇÃO

A crescente urgência em travar as alterações climáticas levou ao desenvolvimento de diversos instrumentos de política ambiental, entre os quais se destacam os mercados de carbono. Estes sistemas procuram internalizar o custo das emissões de gases com efeito de estufa (GEE), criando incentivos económicos para a sua redução. Entre os vários mercados existentes, o Sistema de Comércio de Licenças de Emissão da União Europeia (EU ETS) tornou-se uma das referências internacionais, tanto pela sua escala como pelo grau de maturidade regulatória.

Com o reforço dos compromissos ambientais assumidos pela União Europeia nos últimos anos, e o papel central do EU ETS nas estratégias de descarbonização, o preço do carbono tem vindo a adquirir crescente relevância económica e política. Esta dinâmica levanta um conjunto de questões sobre o seu comportamento, volatilidade e, em particular, a possibilidade de o prever com rigor.

A presente dissertação tem como objetivo analisar a previsibilidade do preço do carbono no mercado EU ETS, com recurso a métodos de *machine learning* e *deep learning*, nomeadamente modelos baseados em redes neuronais profundas, como a arquitetura CNN-LSTM ou mesmo uma arquitetura simples numa Rede Neuronal simples.

Foi construído um conjunto de dados extenso, integrando variáveis macroeconómicas, financeiras, energéticas e climáticas, para testar até que ponto estas características externas podem melhorar o desempenho preditivo. Ao longo do trabalho, foram também considerados modelos mais simples como a *Random Forest*, de forma a permitir uma comparação robusta entre diferentes abordagens.

Os resultados obtidos são analisados à luz da literatura existente e discutidos com base em conceitos como a eficiência dos mercados, a instabilidade estrutural e a presença de mudanças de regime, frequentemente apontadas como barreiras à previsibilidade do preço do carbono.

Questão de Investigação: Será possível prever, numa base diária, os preços referentes às licenças de emissão de carbono, de uma forma confiável e atendendo à elevada incerteza inerente à sua formação?

Relativamente à estrutura adotada, este trabalho encontra-se organizado da seguinte forma: no Capítulo 1 é apresentada a introdução ao tema, enquadrando o problema de investigação e os seus principais objetivos. O Capítulo 2 dedica-se à revisão de literatura, onde são analisados os principais contributos teóricos e empíricos relacionados com os mercados de carbono, metodologias de previsão e as variáveis determinantes no contexto do EU ETS. No Capítulo 3 descreve-se a metodologia adotada, detalhando a selecção de variáveis, o processo de recolha e tratamento de dados, bem como os modelos utilizados. O Capítulo 4 apresenta os resultados empíricos, acompanhados da sua discussão à luz da literatura existente. Por fim, no Capítulo 5, são tecidas as conclusões, salientando as principais contribuições do estudo, as suas limitações e sugestões para investigações futuras.

2. REVISÃO DE LITERATURA

A crescente integração das políticas climáticas na União Europeia, em particular através do Sistema de Comércio de Licenças de Emissão (EU ETS), tem gerado um interesse crescente na compreensão dos mecanismos que influenciam o comportamento do preço do carbono.

A previsão destes preços tornou-se uma área de investigação, não só pela sua importância económica e ambiental, mas também pelos desafios metodológicos que impõe.

Esta secção pretende enquadrar o presente estudo no contexto da literatura existente, explorando os principais contributos académicos sobre a modelação e previsão dos preços do carbono, com especial foco nas abordagens de *machine learning* e *deep learning*. Serão também analisadas as variáveis explicativas habitualmente utilizadas e os desafios associados à instabilidade deste mercado.

2.1 *Sustentabilidade*

2.1.1 *Origem, conceitos e importância*

O conceito moderno de sustentabilidade consolidou-se no final do século XX, tendo adquirido projecção internacional com a publicação do relatório *Our Common Future* da Comissão Brundtland, em 1987. Neste documento, o desenvolvimento sustentável é definido como “aquele que satisfaz as necessidades do presente sem comprometer a capacidade das gerações futuras de satisfazerem as suas próprias necessidades” (*World Commission on Environment and Development*, 1987). Esta definição consagrou o

princípio da equidade intergeracional, que se tornou central no pensamento sobre desenvolvimento e políticas ambientais.

Com o tempo, a sustentabilidade passou a ser entendida como um conceito multidimensional, organizado em torno de três pilares fundamentais: o ambiental, o económico e o social (Purvis, Mao & Robinson, 2019). Esta abordagem tripartida — por vezes representada graficamente pelos círculos interligados “planeta, pessoas e lucro” — expressa a necessidade de conciliar a preservação ecológica, o progresso económico e o bem-estar social. Como salientam Purvis *et al.* (2019), estes três domínios não surgiram simultaneamente nem de forma coordenada, mas foram-se articulando progressivamente em resposta a desafios específicos colocados pelo crescimento económico, pela degradação ambiental e pelas desigualdades sociais.

A evolução do pensamento sobre sustentabilidade foi influenciada por marcos históricos e publicações com grande impacto. Obras como *Silent Spring* (Carson, 1962) e *The Limits to Growth* (Meadows *et al.*, 1972) alertaram para os riscos ecológicos do modelo de crescimento industrial e contribuíram para a consciencialização da opinião pública. Na década de 1980, ganhou força a ideia de que o desenvolvimento económico e a proteção ambiental deveriam ser abordados de forma integrada, culminando na abordagem adotada no relatório Brundtland. A Cimeira da Terra do Rio de Janeiro, em 1992, reforçou esta visão ao colocar o desenvolvimento sustentável no centro da agenda internacional.

Desde então, a sustentabilidade tornou-se um princípio orientador em tratados multilaterais, políticas públicas, estratégias empresariais e até novos paradigmas como a “economia verde” e a “economia circular” (Mohammed *et al.*, 2024). Este princípio está hoje incorporado em iniciativas globais como os Objetivos de Desenvolvimento Sustentável (ODS) das Nações Unidas, adotados em 2015. A definição da Comissão Brundtland está na base destes 17 objetivos, e os três pilares da sustentabilidade estão explicitamente representados nos seus eixos temáticos (UN DESA, 2023).

A importância da sustentabilidade decorre, assim, da sua capacidade de integrar preocupações ambientais, sociais e económicas numa visão coerente de desenvolvimento. Como observa Estoque (2020), apesar da sua ambiguidade conceptual, a sustentabilidade é “de importância fundamental para a humanidade”, funcionando como um ideal normativo — tal como os conceitos de liberdade ou justiça — que orienta ações e projeta visões de futuro.

No entanto, o conceito de sustentabilidade não está isento de críticas. Diversos autores apontam a sua ambiguidade, alertando que expressões como “desenvolvimento sustentável” podem ser demasiado vagas se não forem devidamente contextualizadas (Purvis *et al.*, 2019). Questões como “sustentabilidade de quê e para quem?” revelam a necessidade de uma aplicação concreta e criteriosa do conceito. Como foi notado por Purvis *et al.* (2019), não existe uma teoria unificada subjacente ao modelo dos três pilares, o que dificulta a sua operacionalização rigorosa.

Apesar dessas limitações, a sustentabilidade permanece um conceito mobilizador, com forte poder aspiracional. Ao propor que as necessidades humanas sejam satisfeitas respeitando os limites ecológicos do planeta, oferece um quadro abrangente para orientar decisões em múltiplos domínios — da política à gestão, da tecnologia à educação. É por isso que continua a ser central no discurso académico, nas políticas públicas e nas práticas de governação contemporânea.

2.1.2 O papel das emissões de CO₂ na garantia da sustentabilidade

As emissões de gases com efeito de estufa — em particular o dióxido de carbono (CO₂) — constituem uma ameaça direta à sustentabilidade ambiental e climática. A combustão de combustíveis fósseis e outras atividades humanas têm contribuído para um aumento significativo da concentração de CO₂ na atmosfera, sendo este o principal motor das alterações climáticas antropogénicas (Mackay *et al.*, 2025). A temperatura média global já subiu mais de 1 °C acima dos níveis pré-industriais, e os impactos anteriormente previstos tornaram-se visíveis em fenómenos extremos cada vez mais frequentes, como secas, ondas de calor, inundações e elevação do nível do mar (UNFCCC, 2023).

Como referido por Möller *et al.* (2024), estas alterações colocam em risco o equilíbrio dos ecossistemas e comprometem o bem-estar humano, ameaçando a segurança alimentar, a saúde pública, a infraestrutura e a biodiversidade. A comunidade científica tem vindo a alertar que, sem uma ação célere e eficaz de redução de emissões, o mundo arrisca ultrapassar limites críticos do sistema climático — designados por *tipping points* —, cujos efeitos seriam irreversíveis, como o colapso das calotes polares ou o desaparecimento da floresta amazónica (Möller *et al.*, 2024).

Assim, a sustentabilidade ambiental a longo prazo é incompatível com níveis elevados e contínuos de emissões de CO₂. Como argumenta Fankhauser *et al.* (2022), a descarbonização profunda — ou seja, a redução e posterior eliminação das emissões líquidas de carbono — é condição essencial para garantir a estabilidade climática e, com

ela, as bases ecológicas do desenvolvimento económico e social. O Acordo de Paris (2015), negociado no âmbito da Convenção-Quadro das Nações Unidas sobre Alterações Climáticas, estabeleceu como meta limitar o aquecimento global a bem menos de 2 °C, procurando esforços para não ultrapassar os 1,5 °C (UNFCCC, 2023). Para atingir este objetivo, a comunidade científica recomenda que as emissões globais atinjam o pico até 2025 e sejam reduzidas em cerca de 43% até 2030 (UNFCCC, 2023).

Em termos práticos, isto implica uma transição rápida e estrutural dos sistemas energéticos, com abandono progressivo dos combustíveis fósseis, aumento da eficiência energética, eletrificação dos transportes, promoção de energias renováveis e conservação dos sumidouros de carbono naturais (He *et al.*, 2023). A neutralidade carbónica — definida como o equilíbrio entre emissões e remoção de CO₂ — é hoje amplamente reconhecida como uma condição-chave para travar o aquecimento global e evitar impactos catastróficos (Möller *et al.*, 2024).

Muitos países, cidades e empresas adotaram metas de emissões líquidas nulas para meados do século, inserindo essas metas nos seus planos de sustentabilidade. O conceito de neutralidade carbónica fornece, segundo He *et al.* (2023), um ponto de referência claro para orientar os esforços de descarbonização global. No entanto, a literatura sublinha que estas metas devem ser implementadas com coerência ambiental e alinhadas com os objetivos mais amplos do desenvolvimento sustentável. Fankhauser *et al.* (2022) defendem que a neutralidade carbónica deve ser atingida através de estratégias equitativas e eficazes, assegurando uma transição justa que promova simultaneamente a justiça social, a inclusão económica e a integridade ecológica.

Este princípio está refletido no Objetivo de Desenvolvimento Sustentável 13 da ONU, que apela à ação climática urgente como parte integrante da Agenda 2030 (UN DESA, 2023). Em suma, a redução das emissões de CO₂ não é apenas uma exigência ambiental, mas uma condição de base para garantir a viabilidade dos sistemas naturais e sociais no longo prazo. Cada tonelada de carbono evitada contribui para um futuro mais seguro, resiliente e sustentável.

2.2 Mercados de Licenças de Emissões de CO₂

2.2.1 Propósito e Funcionamento

Os mercados de carbono têm assumido um papel central na arquitetura das políticas de combate às alterações climáticas, funcionando como mecanismos de mercado para limitar as emissões de gases com efeito de estufa (GEE) de forma eficiente e

economicamente sustentável (Montgomery, 1972; Tietenberg, 1974). De uma forma geral, estes mercados dividem-se em dois grandes tipos: os mercados regulados e os mercados voluntários.

Nos mercados regulados, como o Sistema de Comércio de Licenças de Emissão da União Europeia (*EU Emissions Trading System* – EU ETS), aplica-se um modelo de *cap-and-trade*, no qual as autoridades estabelecem um limite absoluto (*cap*) para as emissões totais de determinados sectores económicos. A este teto corresponde uma quantidade equivalente de licenças de emissão (ou *carbon allowances*), que são atribuídas – por leilão ou de forma gratuita – às empresas reguladas. As entidades que emitem menos do que a sua quota podem vender o excedente a outras que ultrapassem os seus limites, criando um mercado secundário de negociação (Ellerman *et al.*, 2010; Schmalensee & Stavins, 2017).

Paralelamente, existem os mercados voluntários de carbono, onde empresas ou indivíduos compram créditos gerados por projetos de redução ou remoção de emissões, geralmente em países em desenvolvimento ou em sectores não cobertos por ETS. Estes créditos são transacionados fora de um enquadramento legal obrigatório e visam compensar emissões residuais, frequentemente como parte de compromissos de neutralidade carbónica corporativa (Bumpus & Liverman, 2008). Embora estes mercados possam desempenhar um papel complementar relevante, a literatura aponta para riscos de integridade ambiental e social, devido à ausência de uniformização regulatória e à dificuldade em garantir a adição e permanência das reduções (Cames *et al.*, 2016; West *et al.*, 2020).

Do ponto de vista económico, os mercados de carbono materializam o princípio da internalização das externalidades ambientais, convertendo as emissões de GEE num custo explícito para os agentes económicos. Ao atribuir um preço ao carbono, estes instrumentos induzem comportamentos de mitigação por parte das empresas, estimulam a inovação tecnológica e promovem a adoção de soluções mais limpas (Stavins, 2008; *World Bank*, 2022). Esta lógica de mercado tem sido validada por diversos estudos empíricos, que demonstram que os sistemas de comércio de emissões – sobretudo quando bem desenhados – conseguem reduzir emissões de forma custo-eficaz, sem impactos negativos relevantes na competitividade das empresas ou no crescimento económico (Martin *et al.*, 2016; Rafaty *et al.*, 2020).

Contudo, a formação dos preços nos mercados de carbono revela-se particularmente complexa e sensível a múltiplos fatores. Em primeiro lugar, os preços são determinados pela interação entre a oferta (controlada pelo *cap*) e a procura, que depende do nível de

atividade económica, da intensidade carbónica dos sectores e dos custos marginais de abate de emissões (Hintermann *et al.*, 2016). Em segundo lugar, os preços reagem a choques exógenos como crises económicas, alterações climáticas extremas ou reformas legislativas. Por exemplo, na Fase I do EU ETS (2005–2007), o excesso de licenças atribuídas levou à queda abrupta do preço para zero, demonstrando a importância de um desenho prudente do mercado (Ellerman & Buchner, 2008).

2.2.2 Principais Mercados Existentes

Os mercados de emissões de carbono, também designados por sistemas de comércio de emissões (em inglês, *Emissions Trading Systems* – ETS), são instrumentos de mercado concebidos para limitar as emissões de gases com efeito de estufa (GEE), atribuindo um preço ao carbono e incentivando a redução de emissões de forma economicamente eficiente (Tietenberg, 2003). Estes sistemas funcionam segundo o princípio do *cap-and-trade*: uma autoridade define um limite máximo (*cap*) de emissões permitido para determinado período, distribui ou leiloa licenças de emissão equivalentes, e permite que estas sejam comercializadas entre os agentes participantes.

A nível mundial, têm vindo a emergir diversos mercados de carbono, com características distintas quanto ao seu desenho, abrangência e grau de maturidade. Entre os mais relevantes destacam-se o Sistema de Comércio de Emissões da União Europeia (EU ETS), o mercado nacional de carbono da China, o sistema da Califórnia (*California Cap-and-Trade Program*), e a Iniciativa Regional de Gases com Efeito de Estufa (RGGI), que abrange vários estados do nordeste dos EUA.

O mercado chinês, operacional desde 2021, é atualmente o maior do mundo em termos de volume potencial de emissões cobertas, tendo começado por abranger exclusivamente o sector da produção de eletricidade a carvão, mas com planos de expansão gradual a outros sectores (Zhang *et al.*, 2022). Já o sistema da Califórnia, criado em 2013, cobre sectores como a eletricidade, a indústria e os transportes, integrando ainda um mecanismo de *offsets* que permite o uso de créditos de emissões provenientes de projetos de mitigação (Burtraw *et al.*, 2018). A RGGI, por sua vez, foi o primeiro mercado de carbono multi-estado nos EUA, iniciado em 2009, e foca-se no sector da eletricidade, utilizando leilões regulares de licenças como principal forma de alocação (Murray & Maniloff, 2015).

Apesar das diferenças, todos estes sistemas partilham o objetivo comum de internalizar o custo ambiental das emissões e criar incentivos económicos para a sua redução. A literatura tem destacado o papel crescente destes mecanismos na governação climática

internacional, embora reconhecendo que o seu impacto varia significativamente consoante a ambição do *cap*, o rigor da monitorização e o funcionamento dos mecanismos de mercado (Meckling *et al.*, 2015).

O EU ETS é amplamente reconhecido como o sistema de comércio de emissões mais desenvolvido e institucionalizado a nível global. Criado em 2005, constitui o principal instrumento da política climática da União Europeia e cobre cerca de 40% das emissões totais da UE (European Commission, 2021). O sistema aplica-se a sectores com elevada intensidade de carbono, como a produção de eletricidade, siderurgia, cimento, refinação e aviação intra-europeia.

A sua evolução tem sido estruturada em quatro fases distintas. A Fase I (2005–2007), considerada piloto, teve como objetivo principal testar a operacionalização do sistema. Esta fase ficou marcada pela alocação excessiva de licenças e pela consequente queda acentuada dos preços, o que revelou fragilidades ao nível do desenho inicial (Ellerman & Buchner, 2008).

A Fase II (2008–2012) coincidiu com o cumprimento do Protocolo de Quioto e introduziu melhorias no sistema de monitorização e verificação. Apesar disso, a crise financeira levou a uma diminuição das emissões e à persistência de um excesso estrutural de licenças, acentuando a volatilidade dos preços (Convery, 2009).

A Fase III (2013–2020) representou uma mudança substancial, com a introdução do leilão como principal método de alocação de licenças, a centralização da governação ao nível da Comissão Europeia, e a criação de um *cap* único em queda anual. Nesta fase, entrou também em funcionamento a *Market Stability Reserve* (MSR), um mecanismo destinado a corrigir os desequilíbrios entre oferta e procura de licenças (Perino & Willner, 2017).

A Fase IV, em curso desde 2021, reforçou a ambição climática da UE no âmbito do Pacto Ecológico Europeu (*European Green Deal*) e do pacote “Fit for 55”. As reformas incluem a redução acelerada do *cap*, o alargamento do sistema a novos sectores, como o transporte marítimo, e a criação de um segundo mercado de carbono para os sectores dos edifícios e transportes rodoviários (European Commission, 2023).

O EU ETS tem sido amplamente analisado na literatura académica, tanto na sua dimensão económica – enquanto mecanismo de precificação do carbono – como na sua eficácia ambiental. Estudos empíricos confirmam que o sistema contribuiu para a redução de emissões, embora com impactos diferenciados por sector e país (Martin *et al.*, 2016). A formação de preços tem sido influenciada por fatores regulatórios, macroeconómicos

e energéticos, com a introdução da MSR a desempenhar um papel importante na recuperação dos preços a partir de 2019 (Fuss *et al.*, 2018).

2.2.3 *Desafio Existentes*

Apesar do seu potencial enquanto instrumentos de política climática, os mercados de emissões enfrentam desafios consideráveis que podem comprometer a sua eficácia. Um dos problemas mais debatidos é a volatilidade dos preços, que dificulta o planeamento de investimento e pode reduzir a previsibilidade dos sinais de mercado. Esta volatilidade resulta frequentemente de choques externos, excesso de licenças ou incertezas regulatórias (Hintermann *et al.*, 2016). Estudos mais recentes confirmam a persistência deste fenómeno, sobretudo em períodos de elevada incerteza energética e geopolítica, como verificado após a invasão da Ucrânia e a crise energética europeia, que intensificaram os desvios nos preços do EU ETS (Chen *et al.*, 2024).

Outro desafio central é o fenómeno de fuga de carbono (*carbon leakage*), que ocorre quando empresas transferem a sua produção para países com regras ambientais menos exigentes, reduzindo a eficácia ambiental do sistema. O EU ETS tem procurado mitigar este risco através da atribuição gratuita de licenças a sectores em risco, embora esta prática levante questões sobre justiça e eficiência (Joltreau & Sommerfeld, 2019). Estudos recentes têm também questionado a eficácia do Mecanismo de Ajustamento de Carbono nas Fronteiras (CBAM) em prevenir esse fenómeno, sublinhando a necessidade de maior coordenação global e de maior transparência nos critérios de aplicação (Nguyen & Branger, 2024).

Adicionalmente, diversos sistemas enfrentaram problemas de sobrealocação de licenças, sobretudo nas fases iniciais, o que conduziu a preços baixos e fraca sinalização de mercado (Ellerman *et al.*, 2010). Apesar das reformas introduzidas — nomeadamente o reforço da *Market Stability Reserve* (MSR) — a literatura recente aponta que continuam a existir desequilíbrios na oferta, com impacto na credibilidade do sinal de preço, especialmente em fases de ajustamento pós-crise (Löw-Beer *et al.*, 2022).

A cobertura limitada de sectores e a escassa integração entre sistemas regionais são também apontadas como limitações à escala e impacto dos mercados de carbono (Schmalensee & Stavins, 2017). Embora se verifiquem esforços para expandir a cobertura setorial e promover ligações internacionais entre mercados (como entre o EU ETS e o mercado suíço), a fragmentação ainda persiste e constitui um obstáculo à eficiência e à uniformização global dos preços de carbono (Meckling *et al.*, 2022).

Por fim, subsistem desafios institucionais e políticos, relacionados com a definição de metas de redução ambiciosas, a credibilidade dos compromissos nacionais e a articulação com outras políticas climáticas e energéticas. A literatura sublinha que, para além de eficiência económica, os mercados de carbono devem garantir integridade ambiental e aceitação social para contribuírem eficazmente para os objectivos do Acordo de Paris (Rafaty *et al.*, 2020). Trabalhos recentes reforçam esta visão, alertando para os riscos de instabilidade política e social associados ao aumento do preço do carbono, especialmente em contextos de inflação energética, e destacam a importância de mecanismos de compensação e justiça climática (Hepburn *et al.*, 2023; Klenert & Mattauch, 2023).

2.3 *Machine Learning*

Nas últimas duas décadas, a área de *Machine Learning* (ML) evoluiu de um domínio de investigação especializado para uma tecnologia transversal com impacto significativo em múltiplos setores. Esta transformação foi impulsionada por três fatores estruturais principais: a crescente disponibilidade de grandes volumes de dados digitais, os avanços substanciais na capacidade computacional e o desenvolvimento de algoritmos mais eficientes e flexíveis (Bontempi, Taieb & Le Borgne, 2023). A conjugação destes elementos permitiu aplicar técnicas de ML a problemas complexos em domínios tão diversos como saúde, finanças, energia, transportes ou ciências sociais (Ahmed *et al.*, 2024).

A importância crescente do ML resulta da sua capacidade para identificar padrões complexos em conjuntos de dados de elevada dimensionalidade, produzindo previsões precisas e permitindo automatizar decisões em contextos incertos (Belkin *et al.*, 2021). Ao contrário das abordagens estatísticas tradicionais, que requerem especificações explícitas das relações funcionais, os métodos de ML aprendem de forma adaptativa a partir dos dados, reduzindo a dependência de pressupostos paramétricos rígidos e oferecendo maior flexibilidade para lidar com sistemas não lineares (Kaufman *et al.*, 2023; Bartlett *et al.*, 2021).

2.3.1 *O que é Machine Learning*

Machine Learning pode ser definido de forma ampla como um conjunto de métodos e técnicas que permitem a um sistema computacional aprender a partir de dados, melhorando o seu desempenho numa tarefa específica sem ser explicitamente programado para tal. Esta aprendizagem ocorre através da identificação de padrões e

regularidades estatísticas em conjuntos de dados, que são depois generalizados para novos casos (Ahmed *et al.*, 2024).

De forma prática, um modelo de ML procura aproximar uma função desconhecida que relaciona um conjunto de variáveis de entrada (*features*) a um resultado de interesse (*target*), utilizando procedimentos estatísticos e computacionais. A escolha do modelo e da função de perda depende do tipo de problema, mas o objetivo subjacente é sempre minimizar o erro preditivo ou maximizar a adequação entre previsões e valores reais (Bontempi *et al.*, 2023).

A definição e a prática de ML não são homogêneas, podendo ser abordadas sob diferentes perspetivas. Numa visão mais operacional, ML é um processo iterativo que envolve (i) a recolha e preparação de dados, (ii) a definição de um modelo e da função objetivo, (iii) o treino e otimização dos parâmetros, (iv) a validação e seleção do modelo mais adequado e (v) a avaliação em dados não utilizados no treino (Kaufman *et al.*, 2023).

Esta perspetiva destaca que ML não é apenas um algoritmo isolado, mas sim um processo metodológico completo, sustentado em princípios estatísticos, computacionais e de engenharia de dados.

2.3.2 Funções Fundamentais em Machine Learning

Independentemente do algoritmo ou do domínio de aplicação, os modelos de *Machine Learning* partilham um conjunto de componentes funcionais que estruturam todo o processo de aprendizagem. Estas funções determinam como o modelo recebe dados, os transforma internamente e produz previsões, bem como a forma como o desempenho é medido e otimizado (Sarker, 2021).

A hipótese representa a função que mapeia as variáveis de entrada (*inputs*) para as variáveis de saída (*outputs*). Pode assumir diferentes formas, dependendo do tipo de problema: modelos lineares, árvores de decisão, *ensemble methods* ou redes neuronais profundas. Esta função pode ser usada para tarefas de regressão, prevendo valores contínuos, ou para tarefas de classificação, atribuindo instâncias a categorias discretas através de funções de ativação adequadas (por exemplo, *softmax* ou sigmoide) (Bontempi *et al.*, 2023).

A função de perda mede a discrepância entre as previsões do modelo e os valores reais. A escolha desta função depende do tipo de tarefa: MSE e MAE são frequentes em regressão, enquanto *cross-entropy loss* e *hinge loss* são amplamente utilizadas em classificação. O objetivo do treino consiste em minimizar esta função, melhorando o

desempenho do modelo sobre os dados de treino e, idealmente, a sua capacidade de generalização (Kaufman *et al.*, 2023).

As funções de ativação introduzem não linearidades nas camadas internas dos modelos, permitindo-lhes representar relações complexas entre variáveis. Funções como ReLU, GELU ou Swish são comuns em arquiteturas modernas, substituindo progressivamente as funções sigmoid e tanh utilizadas em redes mais antigas (Ramachandran *et al.*, 2017).

A regularização controla a complexidade do modelo, evitando que este se ajuste em excesso ao conjunto de treino (*overfitting*). Penalizações L1 e L2, *dropout*, *early stopping* ou *elastic net* são técnicas amplamente utilizadas. Estas abordagens introduzem restrições adicionais no processo de aprendizagem, equilibrando capacidade de representação e generalização (Belkin *et al.*, 2021).

A função objetivo combina a perda com termos de regularização e é minimizada durante o treino. Esta minimização é realizada através de algoritmos de otimização, como *stochastic gradient descent* ou variantes adaptativas (ex.: Adam), que ajustam iterativamente os parâmetros do modelo para reduzir o erro (Sun, Li & Tao, 2020).

$$J(\theta) = \frac{1}{n} \sum_{i=1}^n L(y_i, f_{\theta}(x_i)) + \lambda \Omega(\theta), \quad (1)$$

onde λ controla o peso da penalização (Shalev-Shwartz & Ben-David, 2014).

Na tabela do Anexo I estão resumidas as funções fundamentais do ML. Estas componentes estruturais — modelo, perda, ativação, regularização e otimização — formam a base sobre a qual diferentes algoritmos e arquiteturas são construídos. A escolha e combinação adequadas destes elementos determinam a capacidade preditiva e a eficiência de um modelo em contextos práticos.

2.3.3 Diferenças entre Machine Learning tradicional e Deep Learning

A evolução do ML nas últimas décadas deu origem a duas grandes abordagens: métodos tradicionais e métodos baseados em *deep learning*. Embora partilhem fundamentos comuns, distinguem-se na forma como representam os dados, na escala dos problemas que abordam e na interpretabilidade dos modelos (Ahmad *et al.*, 2024).

Nos métodos tradicionais, a engenharia de variáveis (*feature engineering*) desempenha um papel central. As transformações e seleções de *features* são feitas manualmente, com base em conhecimento especializado sobre o domínio (Sarker, 2021).

Em contraste, o *deep learning* baseia-se na aprendizagem automática de representações através de múltiplas camadas, transformando dados brutos em representações hierárquicas cada vez mais abstratas (Ahmed *et al.*, 2024). Esta capacidade é particularmente útil para dados não estruturados, como imagens, texto ou áudio.

Os métodos tradicionais são eficazes com conjuntos de dados de pequena ou média dimensão e requerem recursos computacionais moderados. Por isso, continuam a ser amplamente utilizados em contextos tabulares e industriais. O *deep learning*, por seu lado, necessita de grandes volumes de dados e de poder computacional elevado (*GPUs* ou *TPUs*), o que lhe permite lidar com problemas de elevada dimensionalidade e grande complexidade (Ahmad *et al.*, 2024).

Os modelos tradicionais têm estruturas relativamente simples, frequentemente lineares ou baseadas em árvores, com um número reduzido de parâmetros. Os modelos de *deep learning* envolvem milhões ou milhares de milhões de parâmetros distribuídos por várias camadas, exigindo técnicas de otimização sofisticadas e longos tempos de treino (Sun *et al.*, 2020).

Os métodos tradicionais são geralmente mais transparentes: modelos lineares ou árvores permitem uma interpretação direta dos coeficientes ou da estrutura. Pelo contrário, as redes profundas são muitas vezes consideradas *black boxes*. Contudo, avanços recentes em *Explainable AI* (XAI), como LIME e SHAP, têm contribuído para aumentar a transparência destes modelos (Vilone & Longo, 2021). No Anexo II estão resumidas as diferenças entre o ML tradicional e o *Deep Learning*.

De forma geral, os métodos tradicionais são mais adequados para conjuntos de dados estruturados, de dimensão reduzida e quando a interpretabilidade é essencial. O *deep learning* é preferido em contextos de grandes volumes de dados não estruturados e tarefas complexas. Em muitos casos, estas abordagens são complementares, sendo frequentemente combinadas em sistemas híbridos que aproveitam as vantagens de ambas (Bontempi *et al.*, 2023).

2.3.4 Como os modelos aprendem e funcionam

Apesar da diversidade de algoritmos e arquiteturas existentes, todos os modelos de *Machine Learning* seguem princípios comuns no processo de aprendizagem. O objetivo é ajustar os parâmetros internos de forma a minimizar o erro entre as previsões e os valores reais, garantindo simultaneamente a capacidade de generalizar para novos dados (Kaufman *et al.*, 2023).

Este processo baseia-se no princípio da minimização do risco empírico. Dado um conjunto de treino $\{(x_i, y_i)\}_{i=1}^n$, procura-se uma função f pertencente a um espaço de hipóteses $\mathcal{F}$ que minimize a perda média observada:

$$\hat{f} = \arg \min_{f \in \mathcal{F}} \frac{1}{n} \sum_{i=1}^n L(y_i, f(x_i)), \quad (2)$$

onde L representa a função de perda. A escolha de L depende da tarefa: funções como MSE e MAE são típicas em regressão, enquanto *cross-entropy loss* é comum em classificação (Bontempi *et al.*, 2023).

$$\theta^{(t+1)} = \theta^{(t)} - \eta \nabla_{\theta} J(\theta^{(t)}), \quad (3)$$

Para encontrar o mínimo da função objetivo, utilizam-se métodos de otimização iterativos, sendo o *stochastic gradient descent* (SGD) e as suas variantes adaptativas (ex.: Adam, RMSProp) os mais frequentes. Estes métodos atualizam os parâmetros do modelo na direção do gradiente negativo da função objetivo, permitindo convergência eficiente mesmo em problemas de grande escala (Sun, Li & Tao, 2020). A introdução de técnicas adaptativas foi determinante para o sucesso do *deep learning*, tornando possível treinar modelos com milhões de parâmetros em tempo razoável.

Outro aspeto crucial é a generalização, isto é, a capacidade do modelo para manter um bom desempenho em dados não vistos. Este desafio é frequentemente conceptualizado através do equilíbrio *bias–variance*. Modelos com elevado *bias* são demasiado rígidos e não capturam padrões relevantes (*underfitting*), enquanto modelos com elevada *variance* ajustam-se excessivamente ao ruído dos dados (*overfitting*) (Bartlett *et al.*, 2021).

A avaliação adequada do desempenho requer procedimentos de validação rigorosos, evitando que o modelo seja avaliado nos mesmos dados em que foi treinado.

A divisão dos dados em conjuntos distintos para treino, validação e teste constitui uma etapa fundamental no desenvolvimento de modelos de ML. Esta separação permite otimizar parâmetros, selecionar modelos e avaliar a capacidade de generalização de forma imparcial, evitando fugas de informação (*data leakage*) entre as diferentes fases (Kaufman *et al.*, 2023).

Treino: nesta fase, o modelo ajusta os seus parâmetros internos com base no *training set*, minimizando a função de perda; Validação: o *validation set* é usado para ajustar hiperparâmetros — parâmetros estruturais como profundidade de árvores, taxas de aprendizagem ou penalizações. Esta afinação evita que as decisões de modelação

dependam dos dados de teste, reduzindo o risco de sobreajuste; Teste: o *test set* é utilizado exclusivamente após a seleção final do modelo e dos hiperparâmetros, fornecendo uma estimativa imparcial do desempenho em dados não observados.

Divisões típicas seguem proporções como 60/20/20 ou 70/15/15, dependendo da dimensão do conjunto de dados. Em situações de amostras reduzidas, recorre-se frequentemente a *k-fold cross-validation*, que divide os dados em *k* subconjuntos e repete o treino *k* vezes, alternando sistematicamente os conjuntos de validação. Esta abordagem melhora a robustez da estimativa do erro fora da amostra (Bontempi *et al.*, 2023).

2.3.5 Tipos de problemas e áreas de aplicabilidade

O *Machine Learning* organiza-se em diferentes paradigmas de aprendizagem, que determinam a forma como os modelos interagem com os dados. Os três mais relevantes são: aprendizagem supervisionada, aprendizagem não supervisionada e aprendizagem por reforço (Ahmed *et al.*, 2024).

Na Aprendizagem Supervisionada os dados estão rotulados, isto é, cada observação possui uma variável de saída associada. O objetivo é aprender uma função que relacione as entradas com as saídas conhecidas e que possa ser aplicada a novos casos. Dentro deste grupo distinguem-se dois tipos principais de problemas: Regressão: previsão de valores contínuos (ex.: preços, consumos energéticos, indicadores financeiros); Classificação: atribuição de instâncias a categorias discretas (ex.: diagnóstico médico, deteção de fraude). Modelos típicos incluem regressão logística, *support vector machines*, *random forests* e redes neurais, avaliados com métricas como MSE, MAE, *F1-score* ou AUC (Bontempi *et al.*, 2023).

Na Aprendizagem Não Supervisionada, os dados não possuem rótulos e o objetivo é descobrir estruturas ocultas. A técnica mais comum é o *clustering*, que agrupa observações semelhantes, utilizando algoritmos como *k-means*, DBSCAN ou *Gaussian Mixture Models*. Outras aplicações incluem redução de dimensionalidade (PCA, *autoencoders*) e deteção de anomalias, úteis em contextos de monitorização e segurança (Sarker, 2021).

Na Aprendizagem Por Reforço, um agente interage com um ambiente e aprende a tomar decisões sequenciais através de recompensas ou penalizações. O objetivo é maximizar a recompensa acumulada ao longo do tempo. O *reinforcement learning* tem aplicações crescentes em robótica, controlo de sistemas complexos, logística e *algorithmic trading* (Lim & Zohren, 2021).

A diversidade destes paradigmas confere ao ML uma enorme versatilidade, tornando-o aplicável em praticamente todos os setores da economia e da ciência. Revisões recentes sublinham a sua relevância em áreas críticas como saúde, finanças, energia sustentável e infraestruturas inteligentes (Raza *et al.*, 2025). Na tabela do Anexo III encontra-se resumido os tipos de aprendizagem e problemas em ML.

2.3.6 Diferenças face aos modelos econométricos tradicionais

A distinção entre ML e modelos econométricos reflete duas culturas de análise de dados com objetivos e pressupostos distintos. A econometria clássica privilegia a inferência causal e a interpretabilidade, baseando-se em modelos paramétricos e em pressupostos estatísticos rigorosos. O ML enfatiza a previsão preditiva fora da amostra, recorrendo a modelos flexíveis e algoritmos de otimização que exigem menos pressupostos formais (Athey & Imbens, 2019). Enquanto a econometria procura explicar relações causais e quantificar efeitos, o ML centra-se em maximizar o desempenho preditivo, mesmo quando a interpretabilidade é limitada.

Os modelos econométricos tradicionais partem de especificações funcionais pré-definidas, frequentemente lineares, e dependem de pressupostos como homoscedasticidade ou ausência de endogeneidade. O ML, em contraste, utiliza modelos não paramétricos e não lineares, ajustando-se aos dados sem a necessidade de pressupostos fortes (Bontempi *et al.*, 2023).

A econometria desenvolveu-se em contextos de dados relativamente pequenos, avaliando modelos com testes estatísticos e diagnóstico de resíduos. O ML foi concebido para lidar com grandes volumes de dados, avaliando modelos com métricas preditivas e procedimentos de validação cruzada (Athey *et al.*, 2021).

Nos últimos anos, têm emergido abordagens híbridas, como o *double machine learning* (DML), que integra técnicas modernas de previsão com a inferência causal, mostrando que os dois paradigmas são complementares (Athey *et al.*, 2021). Na tabela do Anexo IV, as diferenças entre econometria e ML estão resumidas nas diversas dimensões.

2.4 Modelos utilizados na previsão de preços

2.4.1 Features utilizadas na definição dos modelos

Entre as variáveis mais recorrentes nestes estudos encontram-se os preços de energia – nomeadamente do carvão, petróleo, gás natural e eletricidade – dado o seu papel na formação dos custos marginais de abate de emissões (Chevallier, 2011). Adicionalmente,

variáveis macroeconómicas como o índice de produção industrial e o Produto Interno Bruto (PIB) têm sido utilizados como *proxies* da atividade económica e, por conseguinte, da procura por licenças de emissão (Aatola *et al.*, 2013). A temperatura média, sobretudo em períodos sazonais, também tem sido incluída como fator explicativo, dada a sua influência na procura energética para aquecimento ou arrefecimento (Paoletta & Taschini, 2008).

Em modelos GARCH, o foco tem estado na modelação da volatilidade dos preços do carbono. Nestes casos, para além das variáveis anteriores, é comum a inclusão de medidas de volume transacionado, choques de política climática ou efeitos de calendário (Hammoudeh *et al.*, 2014). Ainda que estas abordagens ofereçam uma leitura útil da dinâmica dos preços, têm sido criticadas pela sua limitação na captura de não-linearidades e pela reduzida capacidade de antecipar eventos exógenos.

A aplicação de técnicas de *machine learning* à previsão do preço do carbono permitiu alargar significativamente o leque de variáveis consideradas, beneficiando da maior flexibilidade destes modelos na captura de relações complexas e não-lineares. Abordagens como *Support Vector Machines* (SVM), *Random Forests*, *XGBoost* ou *K-Nearest Neighbours* têm sido utilizadas com sucesso, muitas vezes combinadas com métodos de seleção de variáveis e validação cruzada.

Neste tipo de estudos, para além das variáveis energéticas e macroeconómicas já referidas, tem-se observado a inclusão de índices financeiros globais (como o S&P 500 ou o DAX), preços de *commodities* relacionadas (como cobre, alumínio ou aço), indicadores de política climática, bem como variáveis *proxy* para expectativas de mercado, como volume de pesquisa na internet ou sentimento noticioso (Liu & Lin, 2020; Zhang *et al.*, 2021).

Estas abordagens têm ainda permitido incorporar variáveis geopolíticas, como sanções económicas, conflitos armados ou alterações legislativas internacionais, através da construção de índices compostos ou variáveis *dummy*, revelando-se particularmente úteis em modelos de previsão para horizontes de curto prazo. Apesar do seu potencial preditivo elevado, a literatura alerta para a importância da interpretação destes modelos, sendo a seleção robusta de variáveis um ponto crítico para evitar sobre ajuste e melhorar a generalização (Kumar *et al.*, 2023).

Os modelos de *deep learning*, em especial as redes neuronais recorrentes (RNN) como as LSTM (*Long Short-Term Memory*) e os modelos híbridos CNN-LSTM, têm sido progressivamente aplicados à previsão dos preços do carbono com resultados

promissores. Estas arquiteturas são particularmente eficazes na deteção de padrões temporais complexos e na integração de múltiplas fontes de dados heterogéneos.

Estudos recentes têm demonstrado a eficácia da combinação de variáveis tradicionais com novos dados, incluindo informação climática em tempo real, indicadores de emissões sectoriais, índices de políticas públicas, dados de satélite, consumo de energia por tipo de combustível, e até dados textuais extraídos de relatórios de mercado ou redes sociais (Cheng *et al.*, 2023).

As redes CNN têm sido usadas para extrair automaticamente características relevantes de grandes volumes de dados multivariados (como séries temporais financeiras), que depois alimentam módulos LSTM para modelação da dependência temporal (Shi *et al.*, 2024). Estas arquiteturas híbridas permitem, por exemplo, capturar a reação do preço do carbono a choques macroeconómicos ou geopolíticos com maior sensibilidade do que os modelos tradicionais (Yun *et al.*, 2023).

No entanto, a complexidade destes modelos implica desafios consideráveis em termos de explicabilidade e calibração. A literatura sublinha a necessidade de equilibrar precisão preditiva com transparência, sobretudo quando os resultados são utilizados para orientar decisões de política ou investimento.

2.4.2 Modelos Tradicionais de Previsão

A investigação inicial nesta área assentou maioritariamente em modelos econométricos e de séries temporais clássicos, nomeadamente modelos ARIMA (*AutoRegressive Integrated Moving Average*) e modelos de volatilidade como o GARCH. Estes modelos demonstraram ser úteis na previsão de curto prazo e em contextos de relativa estabilidade (Liu & Huang, 2021), oferecendo ainda a vantagem da na interpretação estatística.

Estudos como os de Benz & Trück (2009) e Chevallier (2011) demonstraram a adequação dos modelos GARCH e de mudanças de regime para captar a heterocedasticidade e os ciclos distintos no EU ETS. Modelos ARMAX (com variáveis exógenas) também foram amplamente utilizados para incorporar fatores macroeconómicos como o PIB, os preços da energia e indicadores industriais (Marfà *et al.*, 2015). Embora estas abordagens tenham permitido compreender relações estruturais importantes, são limitadas na sua capacidade de captar padrões não lineares ou mudanças abruptas de regime – características cada vez mais presentes nos mercados de carbono.

2.4.3 Técnicas de Machine Learning e Deep Learning

Perante as limitações dos modelos tradicionais, a literatura evoluiu no sentido da utilização de técnicas de *machine learning* (ML), nomeadamente redes neurais artificiais (ANN), *Support Vector Machines* (SVR), *random forests* e modelos de *boosting* como o *XGBoost*. Estas técnicas têm a vantagem de não assumirem uma forma funcional prévia entre variáveis, sendo especialmente adequadas para capturar relações complexas e não lineares.

Zhu & Wei (2011) aplicaram SVR na previsão de preços de carbono, tendo reportado melhor desempenho face ao ARIMA. Fan *et al.* (2016) e Zhang *et al.* (2017) utilizaram redes neurais multicamada (MLP) e SVR multivariada, respetivamente, obtendo reduções significativas nos erros de previsão. Mais recentemente, Kumar *et al.* (2024) compararam ARIMA com diversos modelos de *machine learning* e concluíram que *random forests* e *gradient boosting* apresentaram uma performance preditiva substancialmente superior.

Estas abordagens permitiram também a inclusão de variáveis exógenas diversificadas, como indicadores macroeconómicos, preços de energia, emissões por sector, ou até dados alternativos como tendências de pesquisa no Google ou índices noticiosos relacionados com alterações climáticas (Hartvig *et al.*, 2023). A inclusão de tais dados revelou-se particularmente útil para captar expectativas de mercado, choques de informação e efeitos comportamentais não explícitos nos modelos estatísticos clássicos.

A mais recente vaga de investigação tem-se centrado em técnicas de *deep learning*, com particular destaque para as redes neurais recorrentes do tipo LSTM (*Long Short-Term Memory*), que são capazes de modelar dependências temporais de longo prazo. A sua arquitetura com portas de entrada, esquecimento e saída permite-lhes manter memória seletiva ao longo da sequência temporal, sendo particularmente eficazes para séries não estacionárias como os preços de carbono (Zhang & Xia, 2022).

Adicionalmente, as *convulucional neural networks* (CNN), tradicionalmente associadas à análise de imagem, começaram a ser aplicadas a séries temporais. Estas redes são eficazes na deteção de padrões locais e na extração automática de características relevantes, mesmo em dados ruidosos ou com múltiplas variáveis. Shi *et al.* (2024) propuseram um modelo CNN-LSTM para previsão de preços no mercado de carbono de Shenzhen, demonstrando ganhos significativos face a modelos LSTM simples.

Modelos híbridos que combinam CNN com LSTM, ou que integram técnicas de decomposição de séries (como CEEMDAN) com redes profundas, têm demonstrado resultados promissores, nomeadamente em robustez, generalização e velocidade de treino (Yun *et al.*, 2023; Cheng *et al.*, 2024). Estas arquiteturas permitem captar simultaneamente padrões espaciais (entre variáveis) e temporais (ao longo do tempo), o que é particularmente relevante quando se trabalham dados multivariados e de natureza diversa, como é o caso dos mercados de carbono globais.

De uma forma geral, a literatura aponta para uma clara superioridade dos modelos de *machine learning* e *deep learning* relativamente aos métodos estatísticos clássicos no que diz respeito à acuidade preditiva. Os modelos de redes profundas, em particular os híbridos CNN-LSTM, destacam-se pela sua capacidade de aprender padrões complexos e lidar com ruído e não linearidades (Lazcano *et al.*, 2023).

Contudo, esta maior performance vem acompanhada de desafios adicionais, como o risco de *overfitting*, a necessidade de volumes maiores de dados, e a menor capacidade de interpretação dos resultados. Por essa razão, a tendência mais recente tem sido a construção de modelos híbridos, que combinam as vantagens de diferentes métodos e procuram equilibrar performance, explicabilidade e robustez.

A avaliação dos modelos também se tem tornado mais rigorosa, incorporando métricas de estabilidade, previsões multi-horizonte, testes fora da amostra e até robustez a choques estruturais. Este tipo de avaliação é essencial num mercado tão volátil e sensível a eventos externos como o do carbono.

3. METODOLOGIA E DEFINIÇÃO DOS DADOS

De acordo com Fortin (2009), a metodologia corresponde a uma estratégia delineada pelo investigador com o propósito de obter respostas precisas às questões de investigação ou hipóteses formuladas. Pode também ser entendida como o conjunto de métodos e técnicas que orientam o desenvolvimento do processo de investigação científica.

Este capítulo tem como principal objetivo descrever de forma detalhada o processo de recolha e tratamento da informação utilizada na análise. Tal como foi mencionado anteriormente, este estudo procura avaliar se a aplicação de modelos de *machine learning* permite prever com fiabilidade o preço do carbono no mercado EU ETS, contribuindo para uma melhor compreensão da sua dinâmica e possíveis limitações.

3.1 Caracterização do Estudo

O estudo centra-se na análise da previsibilidade do preço do carbono no âmbito do Sistema de Comércio de Licenças de Emissão da União Europeia (EU ETS), com recurso a técnicas de *machine learning* e *deep learning*. A investigação incide sobre dados diários de preços de carbono, recolhidos ao longo das Fases 3 e 4 do mercado, integrando variáveis de natureza macroeconómica, financeira, energética, ambiental e geopolítica.

Trata-se de um estudo de carácter exploratório e explicativo, com o objetivo de perceber não apenas o potencial preditivo dos modelos testados, mas também as limitações impostas pela dinâmica própria deste mercado, frequentemente marcada por episódios de elevada volatilidade e alterações de regime.

3.1.2 Estratégia de investigação adotada

A metodologia deste estudo segue uma abordagem dedutiva, na qual parte de pressupostos teóricos e do conhecimento existente para testar, com base em dados empíricos, a eficácia de diferentes modelos de previsão (Saunders *et al.*, 2015). A investigação assume um carácter quantitativo, na medida em que se recorre à análise estatística e à avaliação rigorosa de desempenho dos modelos aplicados às séries temporais do preço do carbono.

A natureza do estudo é exploratória, uma vez que procura aprofundar o conhecimento sobre um tema complexo e ainda pouco consolidado no contexto nacional, nomeadamente a previsibilidade do preço do carbono no mercado EU ETS através de técnicas de *machine learning*. Este tipo de investigação é particularmente útil quando o objetivo é compreender um fenómeno emergente e propor abordagens metodológicas que possam vir a ser aplicadas em estudos futuros (Robson, 2002).

A estratégia de investigação é baseada em estudos de caso, concentrando-se no comportamento do mercado europeu de carbono (EU ETS), o que se justifica pela complexidade do fenómeno em análise e pela necessidade de compreender em detalhe o funcionamento e os padrões específicos do sistema (Yin, 1984). A amostra utilizada para treino e teste dos modelos é, portanto, não probabilística, uma vez que assenta em dados disponíveis e relevantes para o contexto do caso em estudo (Saunders *et al.*, 2015).

3.1.3 Método de Recolha de Dados

Os dados utilizados foram obtidos através de fontes públicas e institucionais fiáveis, nomeadamente:

- Preços de carbono: recolhidos na *European Energy Exchange* (EEX);

- Indicadores macroeconómicos e financeiros: provenientes de entidades da plataforma *TradingEconomics*;
- Preços de energia, dados ambientais e climáticos: os indicadores foram extraídos do site *Energy-Charts*;
- Contexto geopolítico: O índice GPR foi extraído do site <https://www.matteoiacoviello.com/gpr.htm>

Foi necessária uma harmonização prévia das bases de dados, de modo a garantir coerência na periodicidade e comparabilidade entre as variáveis. Uma das limitações relevantes foi a disponibilidade limitada de alguns dados — como os volumes de emissões e a produção elétrica — que, em muitos casos, só estavam disponíveis a partir de 2015 com a frequência desejada, reduzindo assim a dimensão da amostra.

Esta limitação afetou sobretudo os modelos mais exigentes do ponto de vista computacional, como o CNN-LSTM, cuja performance é altamente sensível ao volume de dados disponíveis para treino.

3.2 Introdução ao caso de estudo

Este caso de estudo tem como principal objetivo avaliar a previsibilidade dos preços das licenças de emissão de carbono no contexto do Sistema de Comércio de Emissões da União Europeia (EU ETS), recorrendo a metodologias de *machine learning* e *deep learning*. Dado o carácter dinâmico e regulado deste mercado, a estratégia metodológica adotada baseou-se numa abordagem exploratória, com iterações sucessivas e reformulação de hipóteses em função dos resultados obtidos. Mais do que desenvolver um modelo otimizado, pretendeu-se compreender os limites estatísticos da previsibilidade e as características estruturais que moldam a evolução dos preços no EU ETS.

4. DEFINIÇÃO DO MODELO

O presente estudo enquadra-se no domínio da aprendizagem supervisionada, mais concretamente no contexto de um problema de regressão, cujo objectivo consiste em prever valores numéricos — neste caso, o preço diário das licenças de emissão de carbono no âmbito do *European Union Emissions Trading System* (EU ETS). Segundo Géron (2023), os problemas de regressão distinguem-se por requererem a estimativa de uma variável-alvo contínua, a partir de um conjunto de preditores (ou *features*) quantitativos e/ou qualitativos, através de um processo de treino em que o modelo aprende padrões históricos entre variáveis independentes e dependentes.

No caso específico deste trabalho, o modelo visa aprender as relações entre os factores macroeconómicos, energéticos e de mercado e o comportamento do preço do carbono, permitindo assim gerar previsões futuras baseadas em dados históricos.

Como salientam Goodfellow, Bengio e Courville (2016), a aprendizagem supervisionada constitui a base da modelação preditiva moderna, sendo amplamente utilizada em finanças, economia e energia para problemas de previsão quantitativa e análise de sensibilidade.

O modelo foi desenvolvido em ambiente *Python*, linguagem amplamente utilizada em investigação científica e *data science* pela sua flexibilidade e vasto ecossistema de bibliotecas *open-source* (VanderPlas, 2023). As principais bibliotecas empregues incluíram *Pandas* e *NumPy* para manipulação de dados, *Scikit-learn* para modelação e avaliação, *Matplotlib* e *Seaborn* para visualização e *TensorFlow/Keras* para o desenvolvimento das redes neuronais (*LSTM* e *CNN*).

Em complemento, foi também utilizado o MATLAB para implementar uma rede neuronal simples (*Feedforward Neural Network*), permitindo comparar o desempenho com abordagens de menor complexidade.

4.1 Escolha das features

A seleção das variáveis explicativas teve por base a literatura económica e empírica sobre os factores determinantes do preço do carbono, com particular incidência nos estudos focados no mercado europeu (Zakeri *et al.*, 2022). As variáveis foram agrupadas em cinco categorias principais: i) preços da energia e dos combustíveis fósseis; ii) indicadores financeiros e de risco; iii) contexto macroeconómico; iv) estrutura de produção energética e emissões; e v) factores geopolíticos.

Os preços do gás natural e do carvão foram seleccionados com base na teoria do *fuel switching*, segundo a qual o diferencial entre combustíveis mais e menos intensivos em carbono influencia diretamente a procura por licenças de emissão (Bunn & Fezzi, 2009). A título complementar, foi incluído o preço do *Brent*, assumido como *proxy* geral da conjuntura económica global.

Do lado dos mercados financeiros, o índice *EuroStoxx 600* e o VIX foram utilizados para capturar o sentimento de mercado e a aversão ao risco (Chevallier, 2011). As taxas de juro das obrigações soberanas alemãs a 2, 5 e 30 anos permitiram captar expectativas económicas e condições monetárias, influenciando indiretamente a procura de licenças (Creti *et al.*, 2012).

Adicionalmente, foram consideradas variáveis relacionadas com o sistema energético e ambiental europeu, nomeadamente a produção de energia por tipo de fonte (fóssil, nuclear, renovável) e as emissões totais de CO₂. Por fim, integraram-se a taxa de juro do BCE, o índice de inflação e o *Geopolitical Risk Index* (GPRD), dados frequentemente referidos na literatura como relevantes para a formação de preços no mercado de carbono (Zakeri *et al.*, 2022). No Anexo V encontram-se resumidas as variáveis e o agrupamento por categoria.

4.2 Pré-Processamento de Dados e Transformações

O pré-processamento dos dados constitui uma fase essencial em qualquer estudo empírico baseado em aprendizagem automática, sendo responsável por assegurar a integridade, a coerência e a comparabilidade das variáveis antes da modelação. De acordo com Fan *et al.* (2021), o pré-processamento “é uma etapa crítica para garantir a fiabilidade dos resultados obtidos a partir de dados operacionais”, permitindo reduzir erros, lidar com falhas e adequar os dados às exigências dos algoritmos. Neste trabalho, o processo de preparação dos dados incluiu duas etapas principais: (i) a imputação de valores em falta (interpolação); (ii) a aplicação de transformações específicas conforme o tipo e natureza das variáveis.

4.2.1 Interpolação dos Dados

O primeiro passo consistiu em tratar as observações em falta, que poderiam comprometer o desempenho dos modelos subsequentes.

Segundo Tawakuli *et al.* (2024), a imputação é uma das operações mais comuns e determinantes em séries temporais, e a escolha do método depende da extensão e natureza das lacunas.

No presente estudo, optou-se por métodos diferenciados conforme o tipo de variável. Na variável *ECB_Rate* foi aplicado o método *forward fill* (preenchimento com o último valor disponível), uma vez que esta variável representa uma taxa de política monetária definida pontualmente pelo Banco Central Europeu, mantendo-se constante até nova decisão. Esta abordagem é consistente com a natureza discreta e persistente das variáveis macroeconómicas (Fan *et al.*, 2021).

Nas restantes variáveis numéricas adoptou-se a interpolação linear entre os valores adjacentes, adequada a falhas curtas e esparsas. Como demonstrado por Mourão Alves de Souza & Lima (2023), a interpolação linear continua a ser um dos métodos mais fiáveis

para lidar com *missing values* em séries contínuas, preservando a tendência local sem distorcer a estrutura temporal.

4.2.2 Transformações

De acordo com Tawakuli, Sánchez e Del Río (2024), a transformação das variáveis é uma das etapas mais críticas do *data preprocessing* em modelos de previsão energética, dado que a maioria dos algoritmos de *machine learning* assume, de forma implícita, uma estrutura estatística relativamente estável e livre de heteroscedasticidade.

Antes da escolha das transformações, foi realizado um diagnóstico da assimetria (*skewness*) de todas as variáveis numéricas, como recomendado por Hyndman e Athanasopoulos (2021). A *skewness* mede a inclinação ou deformação da distribuição em relação à média; valores absolutos elevados ($|skew| > 1$) indicam distribuições com caudas longas, tipicamente associadas a variância instável e presença de *outliers*.

De acordo com Pedregosa *et al.* (2023), a utilização de transformações de potência (tais como logaritmos, *Box-Cox* ou *Yeo-Johnson*) permite reduzir essa assimetria, comprimindo caudas excessivas e aproximando as variáveis de uma distribuição mais simétrica. Assim, optou-se por um método empírico e seletivo, onde cada variável foi transformada apenas quando o seu grau de assimetria o justificava, e sempre com base na natureza do seu domínio (positiva, nula ou negativa).

Cada variável foi avaliada quanto à sua assimetria estatística. Sempre que o valor absoluto de *skewness* ultrapassava 1, considerou-se a necessidade de aplicar uma transformação, de forma a suavizar caudas e estabilizar a variância.

Quando a variável era estritamente positiva e apresentava alta assimetria, aplicou-se o logaritmo natural, definido por:

$$y_t = \ln(x_t) \quad (4)$$

Esta transformação, amplamente utilizada em séries financeiras e energéticas (Hyndman & Athanasopoulos, 2021), converte variações absolutas em variações proporcionais e reduz a heteroscedasticidade, facilitando a identificação de relações multiplicativas. O logaritmo é especialmente útil quando os dados crescem exponencialmente ou variam em várias ordens de magnitude, pois comprime diferenças relativas mantendo a ordem dos valores (Box *et al.*, 2015).

Quando a variável apresentava valores nulos ou negativos, o logaritmo deixava de ser aplicável. Nestes casos, utilizou-se a função inversa do seno hiperbólico (*Inverse Hyperbolic Sine*, *asinh*), definida por:

$$y_t = \operatorname{asinh}(x_t) = \ln(x_t + \sqrt{x_t^2 + 1}) \quad (5)$$

Esta transformação é contínua para todo o domínio dos números reais e comporta-se de forma semelhante ao logaritmo para valores absolutos grandes, mas é aproximadamente linear para valores próximos de zero. Como salientam MacKinnon e Magee (2022), a transformação *asinh* tornou-se uma alternativa amplamente adotada em econometria aplicada, por permitir preservar o sinal, evitar perdas de observações e manter a interpretabilidade económica dos coeficientes. Por outro lado, quando a variável apresentava uma distribuição aproximadamente simétrica ($|skew| \leq 1$), não foi aplicada qualquer transformação, de modo a evitar distorções desnecessárias e manter a integridade da escala original (Pedregosa *et al.*, 2023).

Em Anexo (Anexo VI) encontra-se a tabela resumo das transformações feitas por grupo de variáveis.

4.2.3 Outliers e Correlações

Optou-se por não realizar uma análise específica de *outliers* nem de correlações entre variáveis, decisão metodológica sustentada na natureza dos modelos seleccionados.

Modelos baseados em árvores, como o *Random Forest Regressor*, são notoriamente robustos a *outliers*, uma vez que particionam o espaço de decisão em sub-regiões e limitam o impacto de valores extremos (Breiman, 2001).

Do mesmo modo, redes neurais profundas, incluindo *CNNs* e *LSTMs*, ajustam-se de forma iterativa aos padrões predominantes dos dados, sendo menos sensíveis à multicolinearidade e a observações aberrantes (Shinde & Shinde, 2023).

Assim, em conformidade com as boas práticas actuais de *machine learning*, a análise de *outliers* e correlações não foi considerada uma etapa necessária nesta fase, dado que os modelos a aplicar são intrinsecamente tolerantes a estas condições.

4.2.4 Normalização dos Dados

Concluindo, as variáveis foram normalizadas através do método *Z-score*, utilizando a classe *StandardScaler* da biblioteca *scikit-learn* (Pedregosa *et al.*, 2011). Esta normalização padroniza cada variável de forma a apresentar média nula e desvio-padrão unitário, conforme a expressão:

$$z = \frac{x - \mu}{\sigma} \quad (6)$$

onde μ representa a média da variável e σ o respectivo desvio-padrão.

O método *Z-score* é amplamente recomendado para modelos de *machine learning* baseados em gradiente — como redes neuronais (*CNNs*, *LSTMs*) — pois promove uma convergência mais estável e evita a saturação das funções de activação (Mourão Alves de Souza & Lima, 2023).

Seguindo as boas práticas descritas por Géron (2023), a normalização foi aplicada após o particionamento em conjuntos de treino e teste, ou seja, o *scaler* foi ajustado (*fit*) apenas com os dados de treino e posteriormente aplicado (*transform*) ao conjunto de teste. Esta precaução metodológica impede qualquer forma de *data leakage*, assegurando que a informação do conjunto de teste não influencia o treino do modelo e que a avaliação do desempenho mantém validade estatística.

No fim das transformações foi feita um sumário estatístico. O sumário encontra-se no Anexo VII bem como as figuras que representam as distribuições das variáveis antes e depois das transformações.

4.3 Técnicas de Machine Learning e Deep learning adotadas

A selecção dos modelos de previsão baseou-se na sua relevância científica e na capacidade comprovada de capturar padrões complexos em séries temporais financeiras e energéticas.

Optou-se por explorar três abordagens complementares — CNN-LSTM, *Feedforward Neural Network* (FFN) e *Random Forest Regressor* — representando diferentes paradigmas de aprendizagem supervisionada.

O modelo CNN-LSTM (*Convolutional Long Short-Term Memory*) combina camadas convolucionais com unidades recorrentes do tipo LSTM, permitindo simultaneamente a extração de padrões locais e dependências temporais de longo prazo. Estudos recentes demonstram a eficácia desta arquitetura híbrida na previsão de preços de energia, séries económicas e mercados de carbono, pela sua capacidade de integrar informação sequencial e estrutural de forma hierárquica (Ma *et al.*, 2023). A componente *convolucional* (CNN) identifica padrões espaciais ou locais nos dados — por exemplo, variações sazonais curtas ou choques de mercado — enquanto a LSTM retém memória contextual, essencial em séries com autocorrelação significativa (Goodfellow *et al.*, 2016).

A adoção deste modelo justifica-se, assim, pelo seu destaque na literatura mais recente como uma das abordagens mais robustas e versáteis para previsão em mercados energéticos e ambientais.

Em paralelo, foi implementada uma rede neuronal *feedforward* (*Feedforward Neural Network – FFN*) com uma ou mais camadas ocultas densamente ligadas, usada como análise complementar e linha de base para comparação. Segundo Géron (2023), as FFN constituem o modelo mais elementar de rede neuronal supervisionada, adequado para estabelecer a relação funcional não linear entre variáveis explicativas e dependentes, permitindo avaliar o ganho marginal obtido por arquiteturas mais complexas como a CNN-LSTM.

Por fim, foi incluído o *Random Forest Regressor*, um modelo de ensemble baseado em árvores de decisão introduzido por Breiman (2001), que combina múltiplas árvores treinadas sobre subconjuntos aleatórios dos dados (*bagging*). Este algoritmo destaca-se pela sua robustez a *outliers*, pela capacidade de modelar relações altamente não lineares e pela reduzida necessidade de ajuste de hiperparâmetros. Como salientam Tawkuli *et al.* 2024, os modelos baseados em florestas são amplamente utilizados em problemas de previsão energética e ambiental devido à sua interpretabilidade, estabilidade e resistência à multicolinearidade.

4.4 Métricas de desempenho

A escolha conjunta de MAE, RMSE e R^2 justifica-se por oferecer um equilíbrio entre interpretação intuitiva, penalização de grandes erros e avaliação da variância explicada, conforme recomendado por Géron (2023).

Estas métricas são amplamente reconhecidas na literatura de modelação e previsão de séries temporais por fornecerem perspectivas distintas e complementares sobre a qualidade do ajuste (Hyndman & Athanasopoulos, 2021)

O MAE mede o erro médio absoluto entre os valores previstos ($\hat{y}_i$) e observados (y_i), sendo calculado por:

$$MAE = \frac{1}{n} \sum_{i=1}^n |y_i - \hat{y}_i| \quad (7)$$

Esta métrica expressa a magnitude média dos desvios em unidades da variável dependente, sendo robusta a *outliers* e de fácil interpretação (Chai & Draxler, 2014).

O RMSE corresponde à raiz quadrada da média dos erros ao quadrado, penalizando de forma mais severa grandes desvios:

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2} \quad (8)$$

Por atribuir peso quadrático aos erros, o RMSE é mais sensível a valores extremos, o que o torna particularmente adequado quando se pretende dar maior importância a grandes discrepâncias — uma característica relevante em contextos de previsão económica e energética (Zhang *et al.*, 2022).

Já o coeficiente de determinação (R^2) mede a proporção da variância explicada pelo modelo relativamente à variância total dos dados observados, sendo definido como:

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (9)$$

em que $\bar{y}$ representa a média dos valores observados. O R^2 é uma métrica adimensional e comparável entre modelos, quantificando o grau em que o modelo explica a variabilidade da variável dependente (Hyndman & Athanasopoulos, 2021).

4.5 Estratégias de Análise e Validação

De forma a assegurar uma avaliação robusta e representativa do comportamento do mercado europeu de carbono, foram aplicadas três estratégias complementares de análise e validação: (i) a análise inter-regime e intra-regime, (ii) a seleção de *features* em duas etapas e (iii) a validação temporal com janela deslizante (moving window).

Estas abordagens combinam princípios de estabilidade estatística, redução de dimensionalidade e adaptação dinâmica — aspetos essenciais quando se trabalha com séries financeiras e energéticas sujeitas a regimes regulatórios e estruturais distintos.

4.5.1 Análise Inter-Regime e Intra-Regime

A primeira estratégia consistiu em dividir o conjunto de dados segundo as Fases do EU ETS (*European Union Emissions Trading System*), de forma a avaliar se a previsibilidade do preço do carbono varia ao longo do tempo em função do estágio de maturidade do mercado. O EU ETS, enquanto mercado regulado, sofreu alterações estruturais significativas ao longo das suas fases (I: 2005–2007, II: 2008–2012, III: 2013–2020, IV: 2021–presente), que afectaram profundamente a formação de preços, os volumes transacionados e a liquidez do sistema (Bellassen *et al.*, 2023). Deste modo, a análise inter-regime comparou o desempenho dos modelos entre diferentes fases,

permitindo verificar mudanças no grau de previsibilidade e testar a hipótese de rutura estrutural — isto é, avaliar se as relações entre variáveis explicativas e preço do carbono se mantêm ou alteram consoante o regime de mercado.

Já a análise intra-regime concentrou-se em testar a estabilidade interna de cada fase, verificando se o modelo mantinha capacidade preditiva consistente ao longo de subperíodos homogéneos. Esta abordagem permitiu identificar, empiricamente, se a evolução do EU ETS se traduziu em maior eficiência de mercado (menor previsibilidade) ou se, pelo contrário, a introdução de novos mecanismos (como o *Market Stability Reserve*, na Fase IV) criou padrões interpretáveis. Como observam Zakeri *et al.* (2022), segmentar séries temporais em regimes distintos é particularmente relevante em mercados energéticos regulados, onde as políticas e os mecanismos de ajustamento influenciam diretamente a dinâmica de preços.

4.5.2 Seleção de Features em Duas Etapas

Para reduzir a complexidade e o ruído do modelo, foi implementado um processo de seleção de variáveis (*features*) em duas etapas, com base na importância das variáveis calculada por um *Random Forest Regressor*. Este procedimento segue a lógica de *wrapper selection*, onde a relevância de cada variável é avaliada segundo a sua contribuição para o desempenho do modelo (Guyon & Elisseeff, 2003).

Numa primeira fase, o modelo de *Random Forest* foi treinado sobre um conjunto alargado de *lags* de curto prazo das variáveis explicativas, gerando um ranking de importância através da métrica *Mean Decrease in Accuracy (MDA)*. De seguida, o modelo final foi re-treinado apenas com as variáveis mais relevantes, reduzindo a dimensionalidade sem perda significativa de desempenho.

Este processo garantiu um equilíbrio entre parcimónia e capacidade preditiva, mitigando riscos de *overfitting* e diminuindo o tempo de treino. A seleção foi sempre realizada sem *data leakage*, assegurando que os dados de teste não influenciaram a escolha das variáveis. A metodologia adotada está alinhada com as boas práticas em previsão de séries financeiras e energéticas, onde a eliminação de redundâncias melhora a generalização dos modelos (Tawakuli *et al.*, 2024).

4.5.3 Estratégia de Janela Deslizante (*Moving Window*)

A terceira abordagem metodológica consistiu na implementação de uma janela deslizante temporal (*rolling window* ou *moving window*) como técnica de validação e atualização do modelo.

Esta estratégia foi concebida para refletir a natureza dinâmica e adaptativa dos mercados financeiros e energéticos, nos quais as relações entre variáveis mudam ao longo do tempo (Hyndman & Athanasopoulos, 2021).

Na prática, a janela deslizante consistiu em treinar o modelo com um período móvel de 24 meses, e posteriormente testar o seu desempenho no mês imediatamente seguinte. A cada iteração, a janela avançava um mês, descartando o primeiro ponto do período anterior e incluindo o novo mês mais recente, repetindo o processo de treino e previsão. Este método simula uma estratégia adaptativa, onde o modelo é recalibrado continuamente à luz da informação mais recente — refletindo o comportamento de agentes de mercado que actualizam previsões à medida que novos dados surgem. A aplicação desta técnica no *Random Forest Regressor* permitiu testar a capacidade de generalização e resiliência temporal do modelo, analisando como a sua performance evolui em períodos de maior volatilidade ou estabilidade.

Como referem Bianchi *et al.* (2020), abordagens de janela deslizante são particularmente úteis em mercados sujeitos a *non-stationarity*, pois equilibram a necessidade de memória histórica com a adaptação a novos regimes.

Em anexo, ANEXO VIII, encontra-se o código python dos modelos utilizados para prever os preços.

5. APRESENTAÇÃO E ANÁLISE DOS RESULTADOS

A presente investigação foi conduzida segundo uma estratégia iterativa, com múltiplas fases de testes empíricos. Cada abordagem procurou refinar hipóteses anteriores e explorar diferentes perspetivas sobre a previsibilidade dos preços no mercado do carbono.

Etapa 1 – CNN-LSTM seq2seq com atenção (2015–2025) e Rede Neuronal FeedForward

A primeira abordagem consistiu na aplicação de uma rede CNN-LSTM. O objetivo era prever os log-retornos diários do *Carbon_Price* no período completo de 2015 a 2025, permitindo capturar padrões locais (via convolução) e dependências temporais de longo prazo (via LSTM).

A CNN-LSTM, apesar do seu potencial teórico para captar padrões temporais complexos, revelou um desempenho inferior (RMSE de 4.41€), o que está em linha com o que é reportado por autores como Lim e Zohren (2021), que alertam para o risco de *overfitting* destas redes quando o volume de dados é limitado. Este modelo mostrou-se

particularmente sensível à estrutura dos dados, tendo dificuldades em lidar com um conjunto tabular com alta dimensionalidade e poucos ciclos de treino.

Também foi testado os resultados numa rede neuronal simples, em que o valor do RMSE é aproximadamente 3,90 €. O resultado indica o que já tinha sido visto na CNN-LSTM, que embora com o R^2 bastante positivo, dado a limitação nos dados o modelo pode estar a fazer *overfitting*.

Etapa 2 – *Random Forest Regressor*: avaliação por subperíodos e estabilidade intertemporal

Na segunda fase da análise, recorreu-se ao modelo *Random Forest Regressor*, dada a sua robustez a variáveis colineares, resistência ao sobre ajustamento. O objetivo central foi testar a hipótese de que a imprevisibilidade observada nos retornos diários poderia estar relacionada com mudanças estruturais no mercado, nomeadamente a transição entre a Fase 3 e a Fase 4 do EU ETS. Assim, foram definidos três subconjuntos distintos de treino e teste, de forma a captar possíveis variações na estrutura de dependência temporal entre variáveis.

Em primeiro lugar, testou-se o modelo com a amostra completa, entre 2015 e 2025. O resultado foi um RMSE de aproximadamente 1.13€.

No período correspondente à Fase 3 do EU ETS, caracterizado por menor volatilidade e maior estabilidade regulatória, o modelo apresentou um desempenho muito satisfatório com um RMSE de 0.66€.

No último teste, foi isolada a fase 4 (2021-2025), com maior incerteza e variações abruptas de preço, o erro aumentou para 1.29€. Esta constatação está em consonância com as observações de Chevallier (2012) e Hintermann *et al.* (2016), que destacam a ocorrência de quebras estruturais significativas entre diferentes fases do mercado, com impacto direto na capacidade de generalização dos modelos.

Etapa 3 – Estratégia adaptativa (*Rolling Window*)

A abordagem de janela deslizante revelou-se a mais robusta, ao permitir uma adaptação contínua às dinâmicas mais recentes do mercado. O seu desempenho (RMSE de 1.45€; R^2 de 0.9977) reforça a utilidade de estratégias adaptativas em contextos de elevada variabilidade, como também defendido por Daskalakis *et al.* (2009) e Hammoudeh *et al.* (2014). Ao restringir o treino aos dados mais recentes, o modelo ajusta-se mais eficazmente às alterações de regime e aos choques externos, minimizando o risco

de “contaminação” por padrões desatualizados. A tabela abaixo resume os resultados obtidos.

Tabela I - Tabela com os resultados das métricas RMSE, MAE e R² para os modelos testados

Modelo/ Abordagem	RMSE (€)	MAE (€)	R ²
CNN-LSTM	4.4108	3.7548	0.97
FNN	3.9034	-	0.9922
RandomForest (Dataset Completo)	1.6434	1.1258	0.9959
RandomForest (Fase 3)	0.6597	0.4806	0.9858
RandomForest (Fase 4)	1.2895	1.0598	0.9359
Janela Deslizante (Sliding Window)	1.4461	0.9095	0.9977

A análise comparativa entre os diferentes modelos implementados neste estudo evidencia um conjunto de padrões que merecem ser cuidadosamente interpretados. Em primeiro lugar, a fraca performance do modelo CNN-LSTM — apesar da sua popularidade em tarefas de previsão de séries temporais — sugere que este tipo de rede neural pode não ser o mais adequado para dados com estrutura predominantemente tabular e com um número limitado de observações. O elevado RMSE registado (4.41 €), em contraste com o R² relativamente elevado (0.97), indica que o modelo é capaz de capturar a tendência geral, mas com grande imprecisão nos valores absolutos, o que aponta para um possível problema de *overfitting* ou má adaptação ao tipo de input utilizado (Chevallier, 2011; Lim & Zohren, 2021).

A *Random Forest* revelou um desempenho consideravelmente superior, sobretudo quando os dados foram segmentados por fase do EU ETS. Este resultado reforça a ideia, amplamente referida na literatura (Hintermann *et al.*, 2016; Daskalakis *et al.*, 2009), de que o comportamento do mercado varia significativamente consoante o regime regulatório e o contexto económico de cada fase. De facto, a melhoria substancial na performance quando o modelo é treinado exclusivamente com dados da Fase 3 (RMSE de 0.66 €) sugere que, em contextos mais estáveis e previsíveis, modelos não-paramétricos como o *Random Forest* conseguem capturar bem as dinâmicas do mercado.

Por outro lado, o aumento do erro na Fase 4 (RMSE de 1.29 €), período marcado por maior volatilidade e incerteza regulatória, mostra que a previsibilidade do mercado de carbono diminui em cenários mais dinâmicos e complexos. Este resultado está alinhado com os trabalhos que sublinham a dificuldade em modelar mercados com *structural breaks* frequentes e elevada sensibilidade a fatores exógenos (Byun & Cho, 2013; Creti *et al.*, 2012).

A abordagem baseada em janela deslizante representa a estratégia mais eficaz testada neste estudo, com um R^2 superior (0.9977) e uma redução significativa do erro absoluto (RMSE de 1.45 €), mesmo em períodos de elevada volatilidade. Este resultado destaca a importância de modelos adaptativos em contextos não estacionários, onde a estrutura do mercado está em constante mutação. Ao limitar o treino aos dados mais recentes, a janela deslizante permite ao modelo focar-se na dinâmica atual, evitando a distorção causada por padrões antigos e já ultrapassados. Esta evidência reforça a conclusão de que, em mercados financeiros como o do carbono, a atualidade dos dados é muitas vezes mais valiosa do que a sua quantidade (Zakeri *et al.*, 2022).

Deve ainda ser sublinhado que, apesar do bom desempenho obtido em várias abordagens, os erros absolutos — mesmo nos melhores cenários — não são desprezáveis. Isto confirma a dificuldade inerente à previsão diária do preço do carbono, um ativo influenciado por múltiplos fatores, muitos dos quais não estão diretamente incluídos nos modelos (por exemplo, decisões políticas, expectativas de mercado ou choques energéticos). Como tal, é legítimo afirmar que, embora seja possível melhorar a previsão com técnicas robustas e dados atualizados, persistem limites naturais à precisão atingível em horizontes muito curtos.

6. CONCLUSÕES, LIMITAÇÕES E RECOMENDAÇÕES DE TRABALHO FUTURO

A presente dissertação permitiu verificar, de forma empírica, que a previsão dos preços do carbono no EU ETS é sensível não só ao modelo utilizado, mas também ao regime de mercado em que a previsão é feita e à estrutura temporal dos dados. Os *modelos Random Forest* revelaram-se neste estudo, mais fiáveis do que abordagens *de deep learning* como a CNN-LSTM, principalmente quando aplicados a janelas temporais coerentes e homogêneas.

A principal conclusão deste trabalho prende-se com a importância da adaptação.

Num mercado em constante transformação, métodos que incorporam lógica adaptativa — como a janela deslizante — demonstraram ser mais eficazes do que modelos treinados em longos períodos históricos. A evidência empírica recolhida vem reforçar aquilo que autores como Chevallier (2011) e Convery (2009) já identificavam: os mercados de carbono, apesar de maduros, continuam a apresentar comportamentos instáveis e sensíveis a fatores externos.

Este trabalho enfrentou um conjunto de limitações que importa aqui referir. Em primeiro lugar, a dimensão da amostra, sobretudo no que diz respeito aos dados com frequência diária, constituiu um obstáculo relevante. Muitas das variáveis consideradas relevantes pela literatura — como emissões efetivas de CO₂ ou dados de geração elétrica — só estão disponíveis com o detalhe pretendido a partir de 2015, o que obrigou à exclusão de anos anteriores e reduziu o horizonte temporal útil. Esta restrição comprometeu, em particular, a eficácia de modelos como a CNN-LSTM, cuja performance depende fortemente de grandes volumes de dados para conseguir aprender padrões significativos (Zhang *et al.*, 2021).

Em segundo lugar, não foi possível incluir dados de sentimento de mercado que têm vindo a ganhar destaque em estudos recentes (Zhang & Xia, 2022; Hartvig *et al.*, 2023). A ausência destas fontes de informação pode ter limitado a capacidade dos modelos de capturar certos choques exógenos que influenciam o preço do carbono, como decisões políticas inesperadas, tensões energéticas ou eventos extremos.

Finalmente, a própria estrutura do mercado de carbono, que incorpora mecanismos como a *Market Stability Reserve* (MSR) e ajustes regulatórios frequentes, torna este mercado particularmente desafiante para modelos baseados em relações histórico-estatísticas. Autores como Perino e Willner (2017) demonstram que alterações nas regras do mercado têm impacto imediato nos preços, o que dificulta a aprendizagem de padrões estáveis por parte dos modelos preditivos.

Para investigações futuras, recomenda-se o aprofundamento de abordagens híbridas que combinem *machine learning* com modelos econométricos clássicos (como GARCH), bem como a incorporação de novas fontes de dados (por exemplo, indicadores de sentimento, previsões de volatilidade ou dados climáticos). Também se sugere o uso de previsões intervalares, que poderiam fornecer estimativas mais informativas para apoio à tomada de decisão em contextos de elevada incerteza.

7. REFERÊNCIAS BIBLIOGRÁFICAS

- Aatola, P., Ollikainen, M., Toppinen, A., & Kuosmanen, T. (2013). *Price determination in the EU ETS market: Theory and econometric analysis with market fundamentals*. *Energy Economics*, 36, 380–395.
- Ahmad, T., Zhang, D., Huang, C., & Zhang, H. (2024). *Artificial intelligence in sustainable energy systems: A review*. *Renewable and Sustainable Energy Reviews*, 185, 113615.
- Ahmed, S., Bashar, A., Al Jubair, M., & Moni, M. A. (2024). *Machine learning for healthcare: Review, opportunities and challenges*. *Computer Methods and Programs in Biomedicine*, 249, 108151.
- Agora Energiewende. (2025). *Energy-Charts: Electricity production and market data for EU*. Disponível em <https://www.energy-charts.info/>
- Bartlett, P., Long, P., Lugosi, G., & Tsigler, A. (2021). *Benign overfitting in linear regression*. *Proceedings of the National Academy of Sciences*, 117(48), 30063–30070.
- Belkin, M., Hsu, D., Ma, S., & Mandal, S. (2021). *Reconciling modern machine learning and the bias-variance trade-off*. *PNAS*, 116(32), 15849–15854.
- Bellassen, V., Alberola, E., & Leguet, B. (2023). *Implementing CBAM: The Political Economy of Carbon Border Adjustment in the EU*. *Climate Policy*, 23(4), 497–512.
- Bellemare, M. F., & Wichman, C. J. (2020). *Elasticities and the Inverse Hyperbolic Sine Transformation*. *Oxford Bulletin of Economics and Statistics*, 82(1), 50–61.
- Bontempi, G., Taieb, S. B., & Le Borgne, Y. A. (2023). *Machine learning strategies for time series forecasting*. *European Journal of Operational Research*, 307(1), 1–15.
- Bredin, D., & Muckley, C. (2011). *An emerging equilibrium in the EU emissions trading scheme*. *Energy Economics*, 33(2), 353–362.
- Breiman, L. (2001). *Random Forests*. *Machine Learning*, 45(1), 5–32.
- Bunn, D., & Fezzi, C. (2009). *Structural interactions of European carbon trading and energy prices*. *Journal of Energy Markets*, 2(4), 13–32.
- Chen, C., Liu, H., & Zhang, Y. (2024). *Volatility spillovers between energy markets and the EU ETS: Evidence from the Russia-Ukraine conflict*. *Energy Economics*, 126, 106983.
- Chevallier, J. (2011). *Detecting instability in the volatility of carbon prices*. *Energy Economics*, 33(1), 99–110.

- Chevallier, J. (2012). *A survey of the literature on carbon price forecasting*. *Energy Economics*, 34(3), 807–812.
- Chai, T., & Draxler, R. R. (2014). *Root Mean Square Error (RMSE) or Mean Absolute Error (MAE)? – Arguments against avoiding RMSE in the literature*. *Geoscientific Model Development*, 7(3), 1247–1250.
- Creti, A., Jouvet, P.-A., & Mignon, V. (2012). *Carbon price drivers: Phase I versus Phase II equilibrium?* *Energy Economics*, 34(1), 327–334.
- Daskalakis, G., Psychoyios, D., & Markellos, R. N. (2009). *Modeling CO₂ emission allowance prices and derivatives: Evidence from the European Trading Scheme*. *Journal of Banking & Finance*, 33(7), 1230–1241.
- El Fassed, E. A. A. (2022). *Forecasting volatility in the EU-ETS market: The role of asymmetric GARCH models*. *International Journal of Energy Economics and Policy*, 12(4), 45–53.
- Ellerman, D., Marcantonini, C., & Zaklan, A. (2016). *The European Union Emissions Trading System: Ten years and counting*. *Review of Environmental Economics and Policy*, 10(1), 89–107.
- European Commission. (2021). *EU ETS Handbook (Updated edition)*. Brussels.
- European Energy Exchange (EEX). (2025). *EEX Global Carbon Index – Spot price data*. Disponível em <https://www.eex.com/en/markets/environmentals/global-carbon-indices>
- Fan, C., Chen, M., Wang, X., Wang, J., & Huang, B. (2021). *A review on data preprocessing techniques toward efficient and reliable knowledge discovery from building operational data*. *Frontiers in Energy Research*, 9, 652801.
- Géron, A. (2023). *Hands-On Machine Learning with Scikit-Learn, Keras & TensorFlow* (3rd ed.). O'Reilly Media.
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press.
- Han, Y., Yang, F., & Zhang, Y. (2021). *Research on prediction model of carbon trading price based on BP neural network*. *Environmental Science and Pollution Research*, 28(29), 39001–39015.
- He, B., Jiang, Y., Liu, L., & Zhang, Z. (2023). *Carbon neutrality: A review*. *Journal of Computing and Information Science in Engineering*, 23(6), 1–101.
- Hepburn, C., Klenert, D., & Stern, N. (2023). *Carbon pricing, public support and fairness: Lessons from behavioural economics*. *Nature Climate Change*, 13(1), 1–8.
- Hyndman, R. J., & Athanasopoulos, G. (2021). *Forecasting: Principles and Practice* (3rd ed.). OTexts.

- Iacoviello, M. (2025). *Geopolitical Risk Index*. Recuperado de <https://www.matteoiacoviello.com/gpr.htm>
- Kaufman, S., Rosset, S., & Perlich, C. (2023). *Leakage in data mining: Formulation, detection, and avoidance*. *ACM Computing Surveys*, 55(12), 1–38.
- Klenert, D., & Mattauch, L. (2023). *Political feasibility of carbon pricing and compensation*. *Review of Environmental Economics and Policy*, 17(1), 26–45.
- Lim, J., & Zohren, S. (2021). *Deep learning for financial time series forecasting: A survey*. University of Oxford, Department of Engineering Science.
- Löw-Beer, T., Rosendahl, K. E., & Tietjen, O. (2022). *Is the EU ETS price signal credible? Revisiting market expectations after the MSR reform*. *Energy Policy*, 165, 112959.
- Meckling, J., Nahm, J., & Tørstad, V. (2022). *The rise of climate clubs: How to prevent carbon leakage*. *Foreign Affairs*, 101(6), 138–148.
- Möller, T., Schellnhuber, H. J., Lenton, T. M., et al. (2024). *Achieving net zero GHG emissions critical to limit climate tipping risks*. *Nature Communications*, 15, 6192.
- Nguyen, D., & Branger, F. (2024). *Carbon leakage and competitiveness under the CBAM: Empirical evidence and policy insights*. *Environmental Economics and Policy Studies*.
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., et al. (2011). *Scikit-learn: Machine Learning in Python*. *Journal of Machine Learning Research*, 12, 2825–2830.
- Purvis, B., Mao, Y., & Robinson, D. (2019). *Three pillars of sustainability: In search of conceptual origins*. *Sustainability Science*, 14(3), 681–695.
- Rafaty, R., Dolphin, G., & Pretis, F. (2020). *Carbon pricing and the elasticity of CO₂ emissions*. *Nature Climate Change*, 10, 562–569.
- Sarker, I. H. (2021). *Machine learning: Algorithms, real-world applications and research directions*. *SN Computer Science*, 2(3), 160.
- Shi, H., Sun, Z., & Chen, Y. (2024). *A hybrid CNN-LSTM model for carbon price forecasting: Evidence from the Shenzhen carbon market*. *Applied Energy*, 354, 121450.
- Sutton, R. S., & Barto, A. G. (2018). *Reinforcement Learning: An Introduction* (2nd ed.). MIT Press.
- Sun, R., Li, H., & Tao, D. (2020). *Optimization for deep learning: Theory and algorithms*. *Annual Review of Control, Robotics, and Autonomous Systems*, 3, 239–264.

- Tawakuli, A., et al. (2024). *Time-series data preprocessing: A survey and an insight*. *Journal of Data Science and Analytics*, 6(2), 115–138.
- Trading Economics. (2025). *Global economic indicators database*. Disponível em <https://tradingeconomics.com/>
- UN Department of Economic and Social Affairs. (2023). *Goal 13: Take urgent action to combat climate change and its impacts*.
- UNFCCC. (2023). *The Paris Agreement*. United Nations Framework Convention on Climate Change.
- Vilone, G., & Longo, L. (2021). *Explainable artificial intelligence: A systematic review*. *Information Fusion*, 67, 82–115.
- Yun, Z., Zhang, T., & Tang, B. (2023). *Carbon price prediction with CEEMDAN and LSTM: A hybrid approach*. *Energy Economics*, 126, 106871.
- Zakeri, B., Syri, S., & Connolly, D. (2022). *Modeling the future of carbon pricing: Empirical analysis using machine learning*. *Applied Energy*, 306, 118020.
- Zhang, B., & Xia, Y. (2021). *Big data and carbon price forecasting: Integrating online search data and news sentiment*. *Resources Policy*, 74, 102258.
- Zhang, X., Wang, C., & Song, D. (2023). *Extreme events, carbon market volatility and risk transmission: Empirical evidence from the EU ETS*. *Environmental Science and Pollution Research*, 30, 67212–67229.

8. ANEXOS

Anexo I – Funções fundamentais em *Machine Learning*

Tabela II - Tabela com o resumo das funções fundamentais do *Machine Learning*

Tipo de Função	Papel Principal	Exemplos
Hipótese / Modelo	Mapear inputs → outputs para regressão ou classificação	Regressão linear, SVM, <i>random forests</i> , redes neuronais (ReLU, softmax)
Perda (Loss)	Medir discrepância previsão-realidade	MSE, MAE, cross-entropy, hinge
Otimização / Objetivo	Minimizar perda + regularização	$J(\theta) = L(\theta) + \lambda\Omega(\theta)$
Ativação	Introduzir não linearidades internas	Sigmoid, tanh, ReLU, softmax, Swish, GELU
Regularização	Controlar complexidade / evitar <i>overfitting</i>	L1, L2, elastic net, dropout, early stopping

Anexo II – *Machine Learning* vs *Deep Learning*

Tabela III - Tabela que resume as principais diferenças entre o ML tradicional e o DL

Dimensão	ML Tradicional	Deep Learning
Representação	<i>Feature engineering</i> manual	Aprendizagem automática de representações
Dados necessários	Pequenos a médios conjuntos	Grandes volumes (big data)
Custo computacional	Baixo / moderado	Elevado (GPU/TPU)
Arquitetura	Estruturas simples, poucos parâmetros	Redes profundas, milhões de parâmetros
Interpretabilidade	Alta (modelos transparentes)	Baixa (black box), mitigada com XAI
Domínios típicos	Dados tabulares, económicos, industriais	Imagens, texto, áudio, problemas complexos

Anexo III – Tipos de aprendizagem e classificação dos problemas em *Machine Learning*

Tabela IV- Tabela com os tipos de aprendizagem e as categorias de problemas em ML

Aprendizagem	Tipo de Problema	Objetivo Principal	Exemplo Prático
<i>Supervised</i>	Regressão	Prever valores contínuos	Previsão de preços, consumo energético
<i>Supervised</i>	Classificação	Atribuir categorias discretas	Deteção de fraude, diagnóstico médico
<i>Unsupervised</i>	Clustering	Agrupar dados sem rótulos	Segmentação de clientes, padrões ocultos
<i>Unsupervised / Hybrid</i>	Dimensionalidade / Anomalias	Reduzir complexidade ou detetar desvios	PCA, autoencoders, deteção de fraude
<i>Reinforcement</i>	Decisão sequencial	Aprender políticas ótimas com base em recompensas	Robótica, controlo, trading algorítmico

Anexo IV – Diferenças fundamentais entre os modelos tradicionais econométricos e o *Machine Learning*

Tabela V - Tabela que resume as diferenças entre modelos econométricos e modelos de ML

Dimensão	Econometria	Machine Learning
Objetivo principal	Inferência causal, interpretação	Previsão fora da amostra, desempenho preditivo
Estrutura funcional	Paramétrica, linear ou especificada	Flexível, não paramétrica, não linear
Pressupostos	Fortes (exogeneidade, homoscedasticidade)	Mínimos, foco em padrões empíricos
Dados	Pequenos, baixa dimensionalidade	Grandes volumes, alta dimensionalidade
Avaliação	Testes estatísticos, diagnóstico de resíduos	Métricas preditivas, validação cruzada
Interpretabilidade	Alta, coeficientes com significado económico	Baixa (sobretudo DL), mitigada com XAI
Abordagens híbridas	Tradicionalmente limitadas	DML, causal forests, integração de métodos

Anexo V – Variáveis escolhidas com base na literatura, para a construção do modelo de previsão*Tabela VI - Tabela com a descrição de cada variável e a que categoria pertencem*

Variável	Descrição	Categoria
Brent Price	Preço do Barril de Petróleo - Brent	i) Preços da energia e dos combustíveis fósseis
Coal Price	Preço do carvão	i) Preços da energia e dos combustíveis fósseis
Gas Price	Preço do gás natural - TTF	i) Preços da energia e dos combustíveis fósseis
EDPR Price	Cotação da EDP Renováveis	ii) Indicadores financeiros e de risco
EuroStoxx Price	Índice das maiores empresas Europeias	ii) Indicadores financeiros e de risco
Germany 2Y Bond	Yield da obrigação alemã a 2 anos	ii) Indicadores financeiros e de risco
Germany 30Y Bond	Yield da obrigação alemã a 30 anos	ii) Indicadores financeiros e de risco
Germany 5Y Bond	Yield da obrigação alemã a 5 anos	ii) Indicadores financeiros e de risco
VIX Price	Índice de volatilidade implícita (VIX)	ii) Indicadores financeiros e de risco
ECB Rate	Taxa diretora do BCE	iii) Contexto macroeconómico
Inflation	Taxa de inflação da Zona Euro	iii) Contexto macroeconómico
Co2 Mt	Emissões totais de CO ₂	iv) Estrutura de produção energética e emissões
EU avg spot price	Preço médio da eletricidade na UE	iv) Estrutura de produção energética e emissões
Fossil MW	Produção fóssil total	iv) Estrutura de produção energética e emissões
Nuclear Energy Index	Índice de produção nuclear ponderada	iv) Estrutura de produção energética e emissões
Nuclear MW	Produção nuclear total	iv) Estrutura de produção energética e emissões
Renewable MW	Produção renovável total	iv) Estrutura de produção energética e emissões
Solar Energy Index	Índice de produção solar ponderada	iv) Estrutura de produção energética e emissões
EU_ETS_Fase	Fase do sistema de comércio de emissões (III ou IV)	v) Fatores geopolíticos
GPRD	Geopolitical Risk Index (risco político global)	v) Fatores geopolíticos
Carbon Price	Preço do EU ETS	Variável Alvo

Anexo VI – Diagnostico de assimetria e transformações aplicadas por variável*Tabela VII - Tabela com a assimetria por variável e transformação aplicada em função do resultado do teste*

Variável	Skewness	Transformação
EU avg spot price	2.698360541	log
Gas Price	2.621258921	log
VIX Price	2.530838604	log
EuroStoxx Price	2.443986734	log
Coal Price	2.102959067	log
GPRD	1.763317762	log
Inflation	1.584509282	asinh
ECB Rate	1.321963386	asinh
Germany 2Y Bond	1.060549771	asinh
Solar Energy Index	0.894917463	NA
Germany 5Y Bond	0.892949188	NA
EDPR Price	0.630840829	NA
Nuclear Energy Index	0.597830824	NA
Renewable MW	0.541892234	NA
Co2 Mt	0.519807664	NA
Carbon Price	0.443767028	NA
Germany 30Y Bond	0.409807129	NA
Brent Price	0.164968585	NA
Nuclear MW	0.079900089	NA
Fossil MW	-0.011584332	NA

Anexo VII – Sumário estatístico dos dados já tratados e pronto a utilizar no modelo.
Distribuições de cada variável antes e depois das transformações.

Tabela VIII - Tabela com o sumário estatístico das variáveis que irão servir como base do modelo preditivo

Variável	n	missing	mean	std	min	q25	median	q75	max	skew	kurtosis
Brent Price	2727	0	5.00E-16	1	-2.676979813	-0.793140821	-0.009561854	0.684229053	3.260433746	0.164968585	-0.122322566
Carbon Price	2727	0	-1.67E-16	1	-1.118896306	-0.988801079	-0.424843043	0.970897433	2.100403911	0.443767028	-1.362844625
Co2 Mt	2727	0	-2.50E-16	1	-2.03750999	-0.736449975	-0.145212594	0.573143817	3.136290274	0.519807664	-0.110302761
Coal Price	2727	0	-4.17E-17	1	-1.398543893	-0.804922372	-0.128568893	0.512639341	2.669272971	0.89223447	0.34169955
ECB Rate	2727	0	1.67E-16	1	-0.605349816	-0.605349816	-0.605349816	1.072974949	1.963167678	1.158470644	-0.542840939
EDPR Price	2727	0	-4.17E-16	1	-1.191411317	-0.902488079	-0.422414105	0.900937141	2.400059178	0.630840829	-1.046443833
EU avg spot price	2727	0	-1.25E-16	1	-3.453456197	-0.717111664	-0.266539934	0.605652755	3.487650998	0.86348012	0.466172431
EuroStoxx Price	2727	0	1.06E-15	1	-1.954092077	-0.766365009	-0.091865463	0.617373616	4.889986998	0.717945152	0.863151628
Fossil MW	2727	0	4.17E-17	1	-2.773175588	-0.650865625	0.007865404	0.688586946	2.955763215	-0.011584332	-0.24349697
Gas Price	2727	0	-5.00E-16	1	-2.788694512	-0.639023934	-0.20733069	0.554927519	3.263519757	0.543601685	0.321439729
Germany 2Y Bond	2727	0	-1.25E-16	1	-0.969364354	-0.714937646	-0.618288856	1.295247398	1.891076591	0.935527817	-0.957953581
Germany 30Y Bond	2727	0	-6.25E-17	1	-1.769157439	-0.880885494	-0.096430664	0.817795559	2.122395321	0.409807129	-0.94613124
Germany 5Y Bond	2727	0	4.17E-17	1	-1.129356701	-0.740707153	-0.501411254	1.291657678	2.093711522	0.892949188	-0.938257474
GRPD	2727	0	1.17E-15	1	-5.605182311	-0.610938075	0.023565096	0.663432187	3.667753844	-0.251833887	0.914401572
Inflation	2727	0	8.34E-17	1	-2.203424633	-0.68547425	-0.048240162	0.425694036	2.012839357	0.234174217	-0.558789595
Nuclear Energy Index	2727	0	2.08E-16	1	-1.749251956	-0.857711622	-0.355924319	0.687493461	2.875712148	0.597830824	-0.805896781
Nuclear MW	2727	0	-4.17E-17	1	-4.07394713	-0.687300616	0.020423569	0.574333412	2.549778226	0.079900089	-0.230429387
Renewable MW	2727	0	2.08E-16	1	-2.150246103	-0.72894078	-0.126269891	0.646355268	3.726635025	0.541892234	-0.044974497
Solar Energy Index	2727	0	1.67E-16	1	-1.126273836	-0.809656316	-0.39788949	0.945069147	3.195227206	0.894917463	-0.312189571
VIX Price	2727	0	-2.50E-16	1	-1.924217713	-0.764805257	-0.144564827	0.640258826	4.711258315	0.730951102	0.811782787

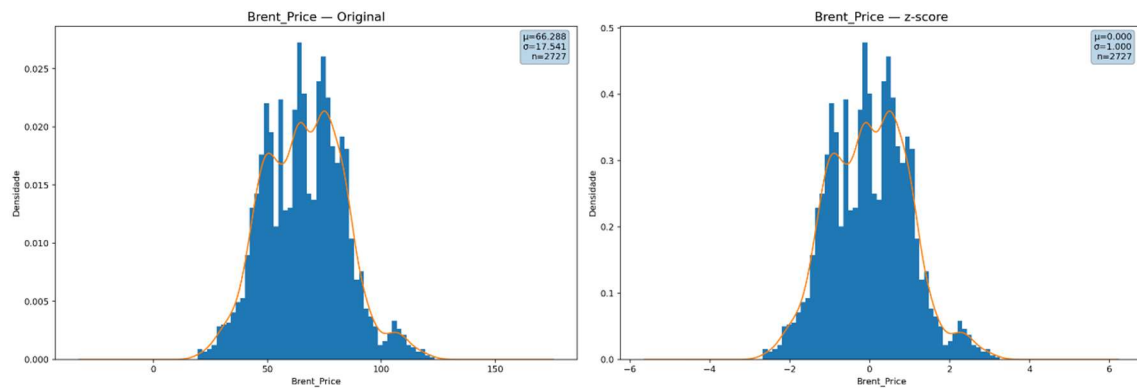


Figura I - Distribuição Brent_Price antes vs depois

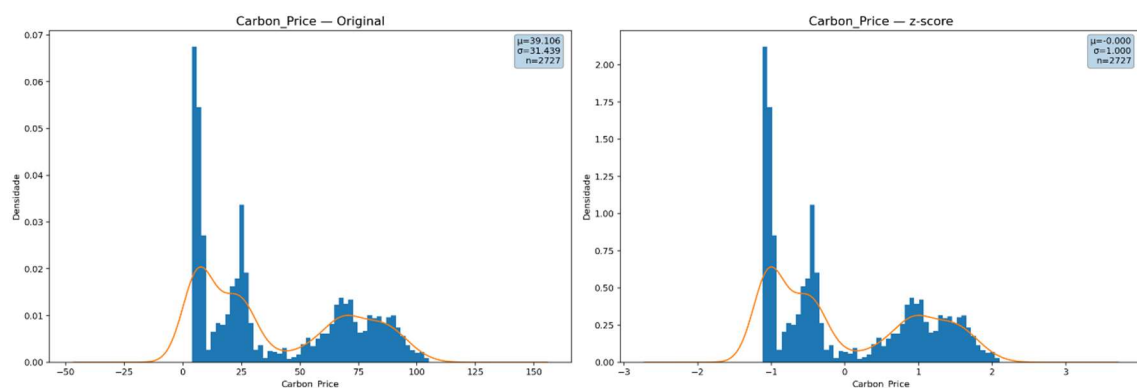


Figura II - Distribuição Carbon_Price antes vs depois

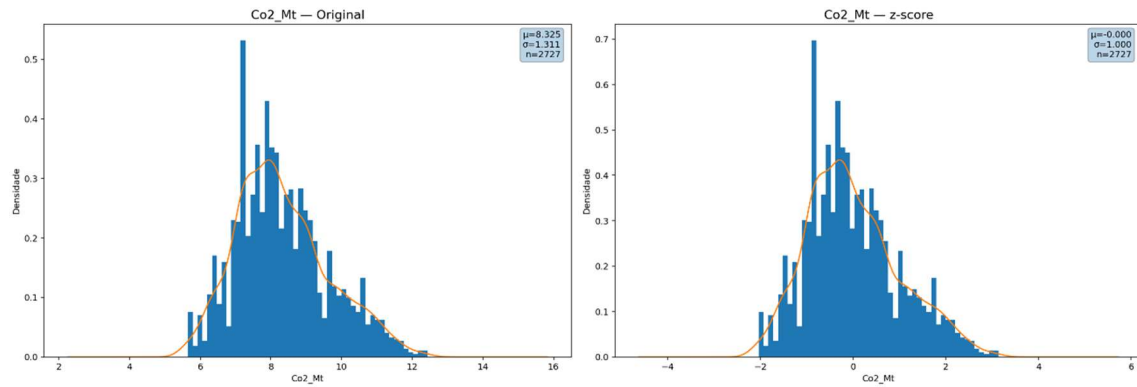


Figura III - Distribuição CO2_Mt antes vs depois

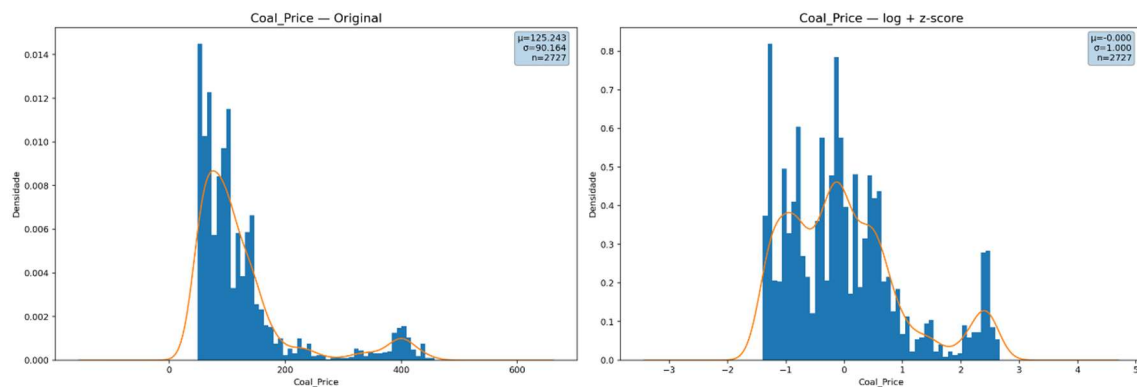


Figura IV - Distribuição Coal_Price antes vs depois

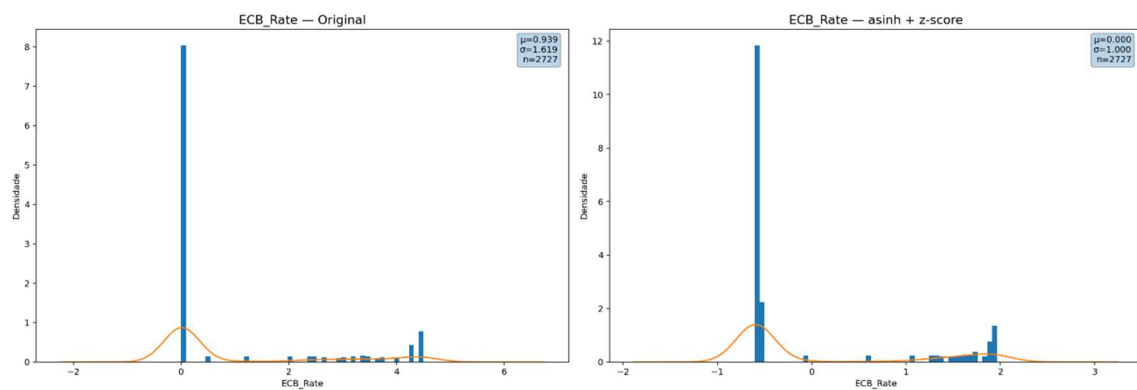


Figura V - Distribuição ECB_Rate antes vs depois

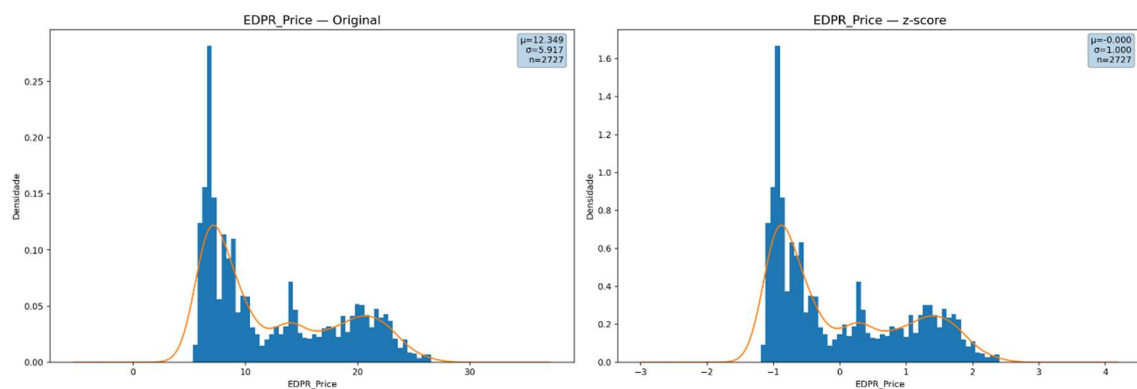


Figura VI - Distribuição EDPR_Price antes vs depois

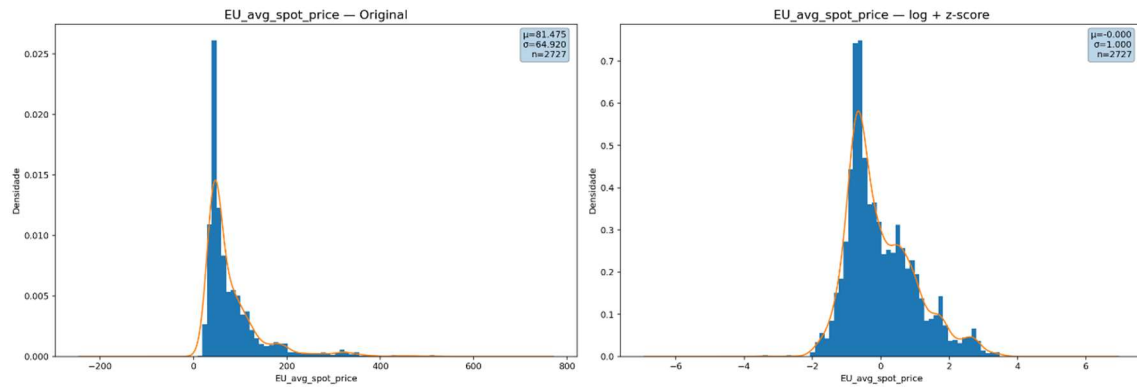


Figura VII - Distribuição $EU_avg_spot_price$ antes vs depois

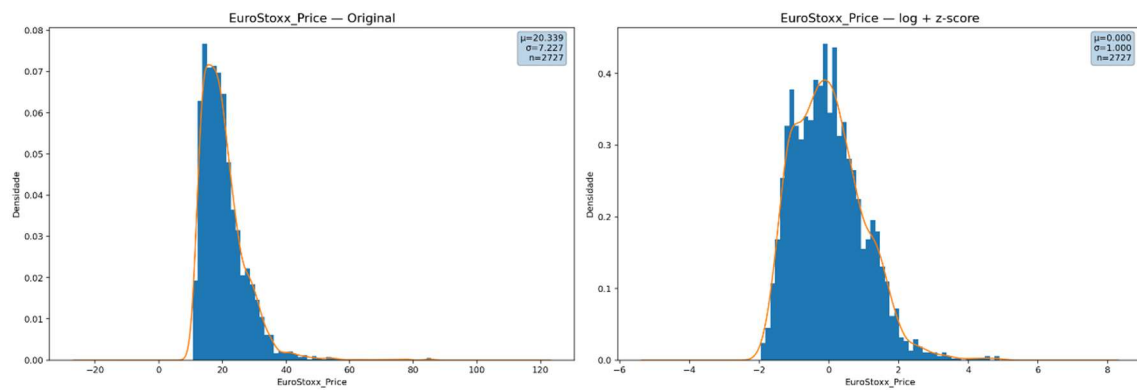


Figura VIII - Distribuição $EuroStoxx_Price$ antes vs depois

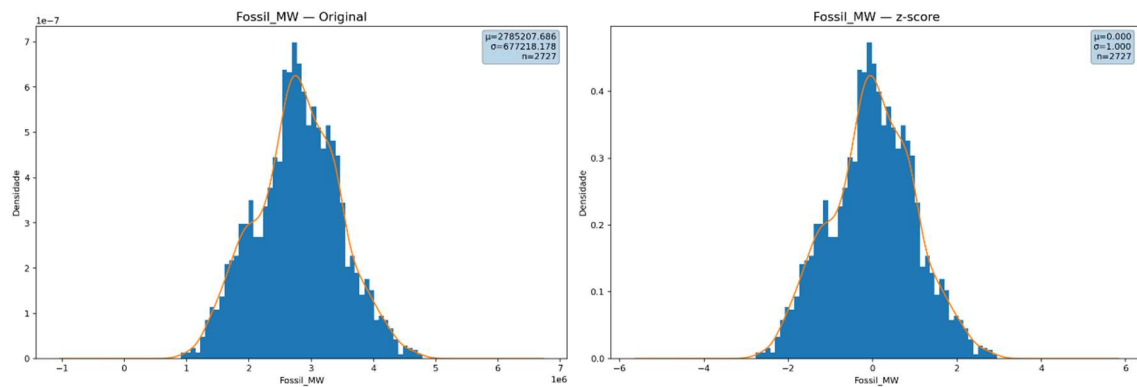


Figura IX - Distribuição $Fossil_MW$ antes vs depois

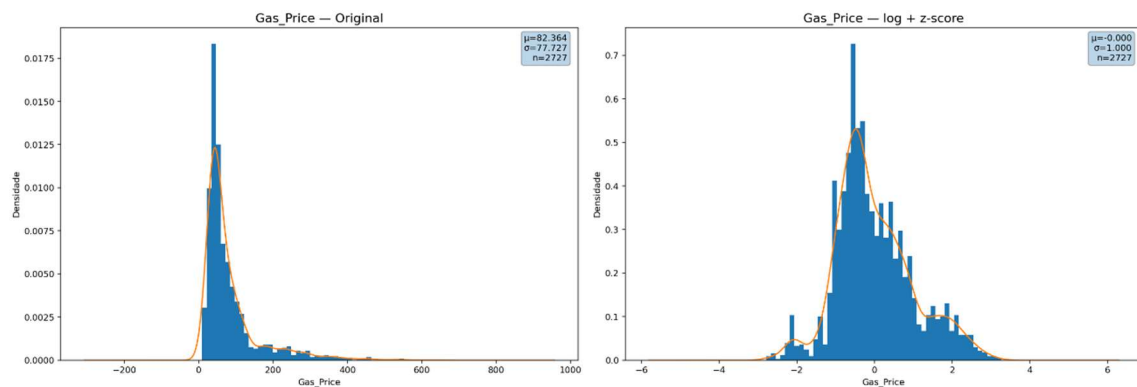


Figura X - Distribuição Gas_Price antes vs depois

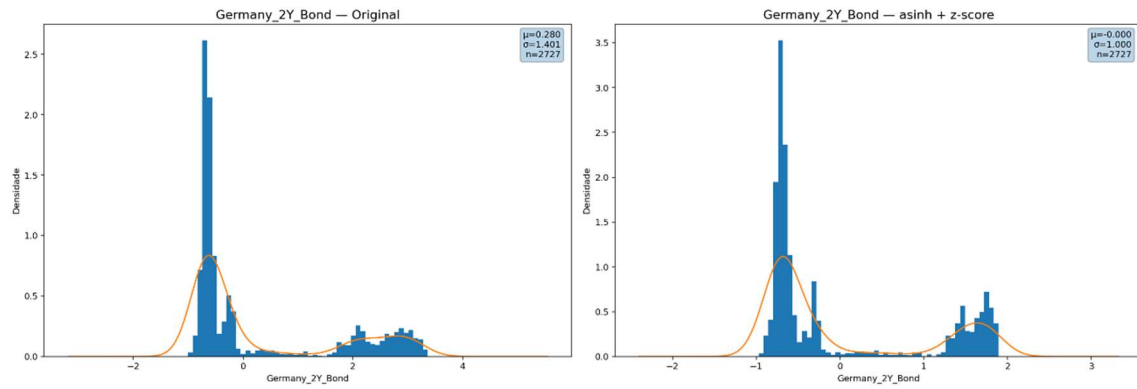


Figura XI - Distribuição German_2Y_Bond antes vs depois

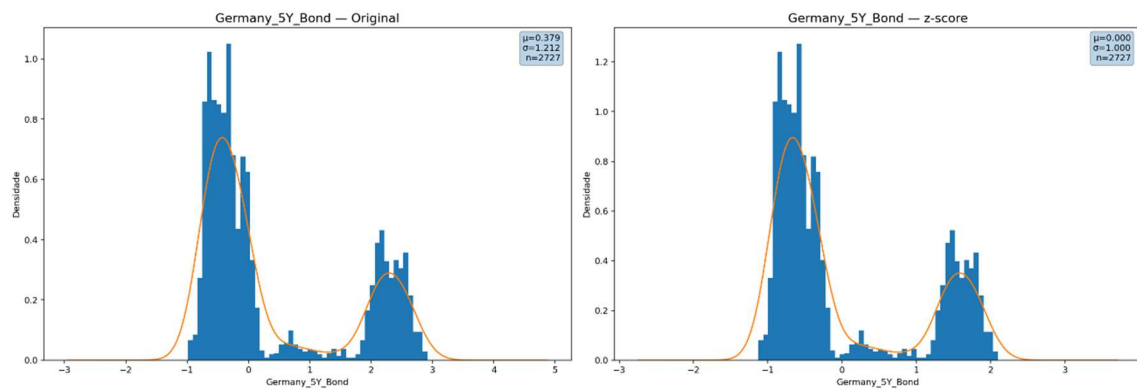


Figura XII - Distribuição German_5Y_Bond antes vs depois

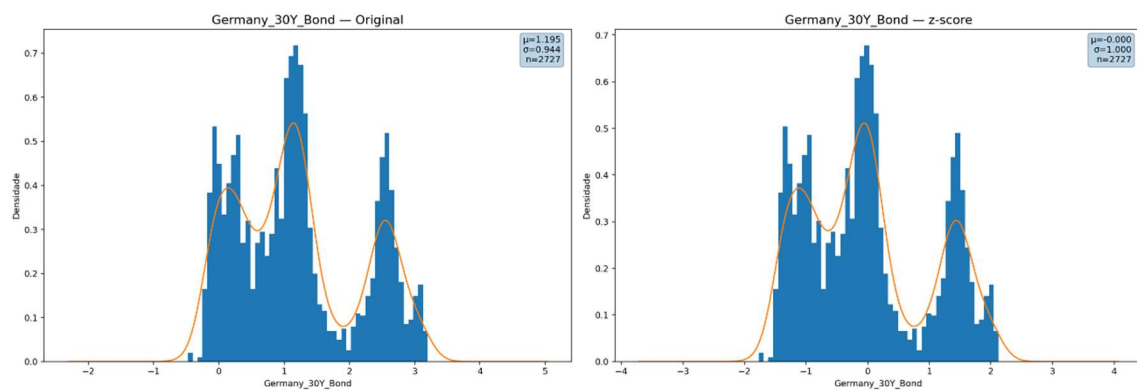


Figura XIII - Distribuição German_30Y_Bond antes vs depois

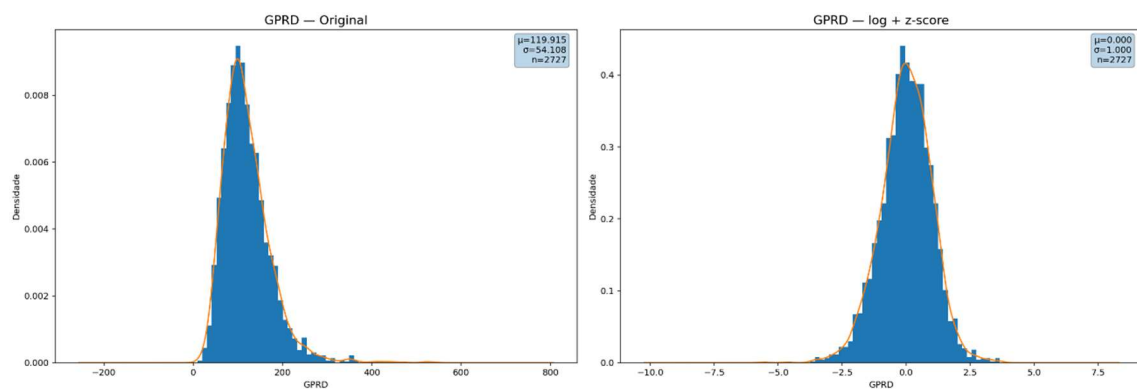


Figura XIV - Distribuição GPRD antes vs depois

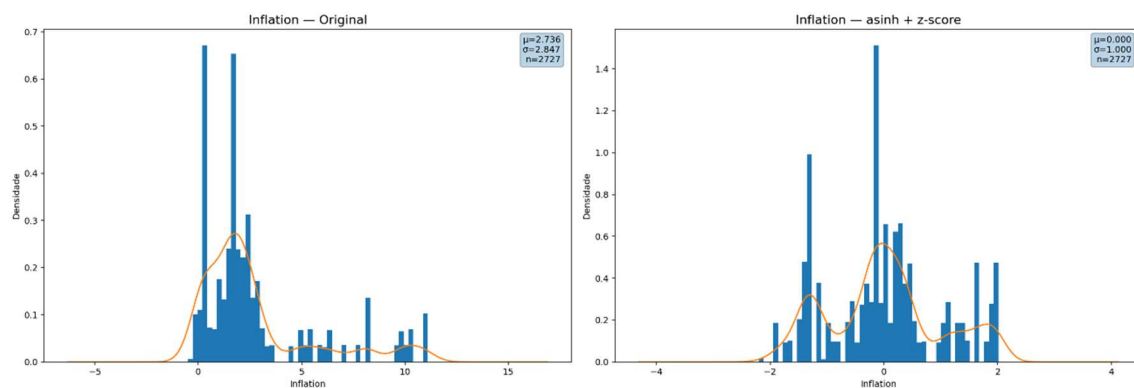


Figura XV - Distribuição Inflation antes vs depois

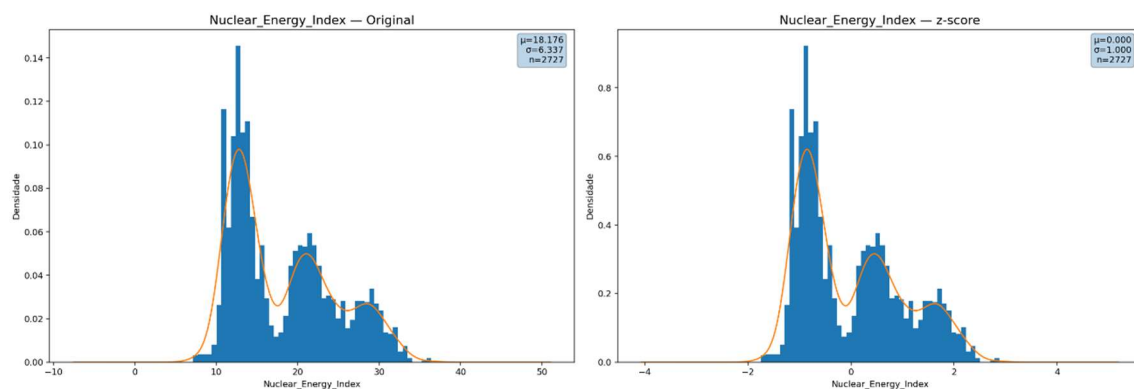


Figura XVI - Distribuição Nuclear_Energy_Index antes vs depois

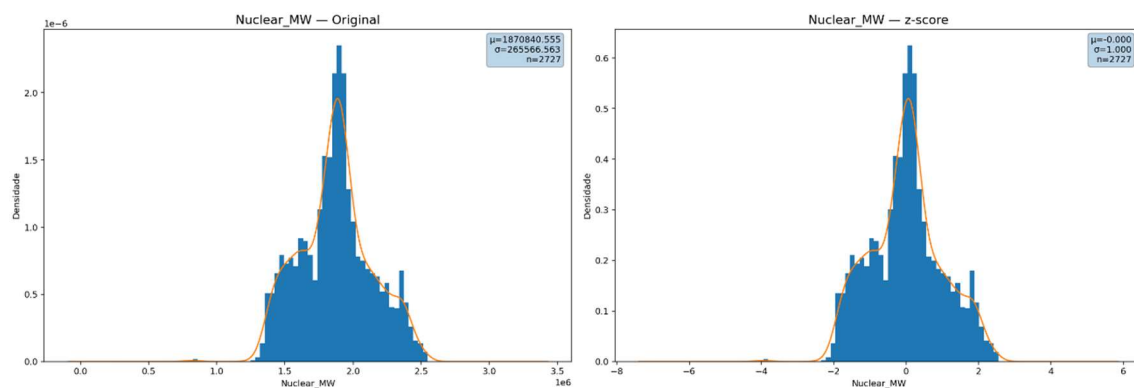


Figura XVII - Distribuição Nuclear_MW antes vs depois

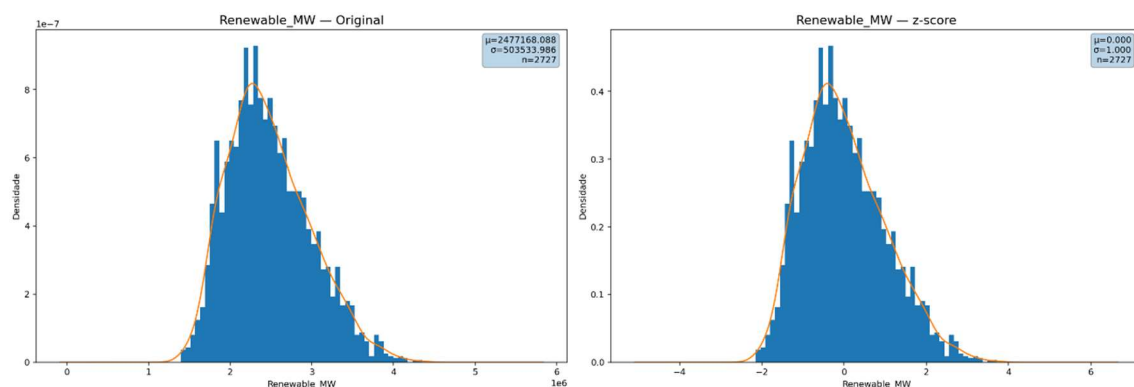


Figura XVIII - Distribuição Renewable_MW antes vs depois

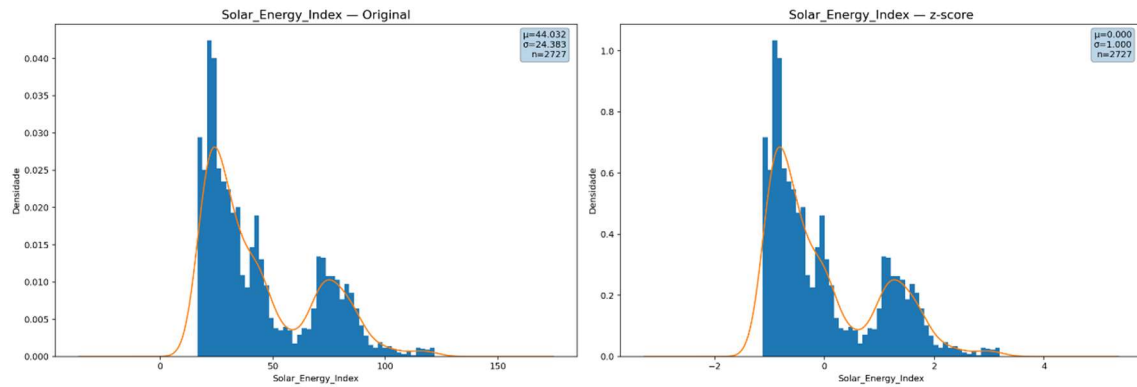


Figura XIX - Distribuição Solar_Energy_Index antes vs depois

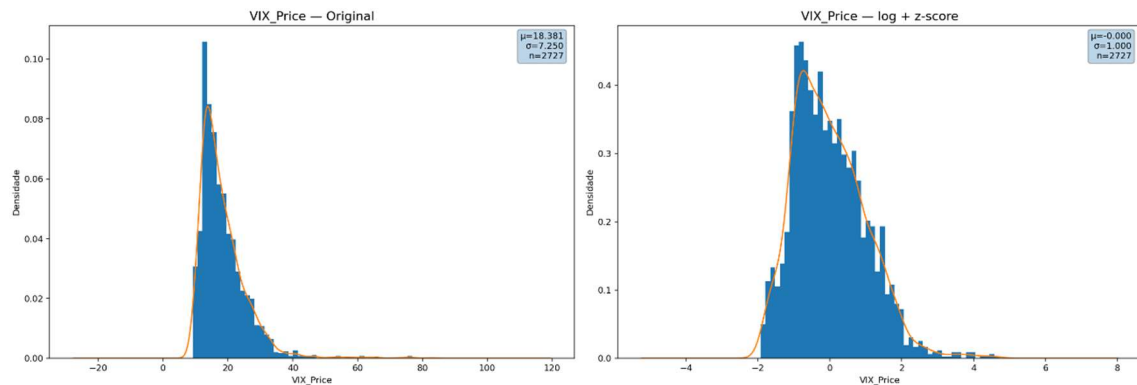


Figura XX - Distribuição VIX_Price antes vs depois

Anexo VIII – Código Python dos modelos usados

```
# Definir arquitetura CNN-LSTM
def build_light_cnn_lstm(input_shape, forecast_horizon=1):
    model = Sequential()
    model.add(Conv1D(filters=8, kernel_size=3, activation='relu', padding='same', input_shape=input_shape))
    model.add(Dropout(0.3))
    model.add(LSTM(32))
    model.add(Dropout(0.3))
    model.add(Dense(forecast_horizon))
    model.compile(optimizer='adam', loss='mse', metrics=['mae'])
    return model

# Dados de entrada
input_shape = (X_train.shape[1], X_train.shape[2])
forecast_horizon = y_train.shape[1]

# Construir e treinar o modelo
model = build_light_cnn_lstm(input_shape, forecast_horizon)

early_stop = EarlyStopping(monitor='val_loss', patience=10, restore_best_weights=True)

history = model.fit(
    X_train, y_train,
    epochs=100,
    batch_size=32,
    validation_split=0.2,
    callbacks=[early_stop],
    verbose=1
)

# Previsões
y_pred = model.predict(X_test)

# Achatar arrays para 1D (se forecast_horizon = 1)
y_test_flat = y_test.reshape(-1)
y_pred_flat = y_pred.reshape(-1)
```

Figura XXI- Estrutura da CNN-LSTM usada para testar a previsão do preço do EU ETS

```

#Seleção de Features
ranking_features_cols = []
for feature in X_train.columns:
    if 'lag_0' in feature or (int(feature.split('_')[1]) <= N_RANKING_LAGS):
        ranking_features_cols.append(feature)

X_ranking_train = X_train[ranking_features_cols]

ranking_model = RandomForestRegressor(n_estimators=100, random_state=42, n_jobs=-1)
ranking_model.fit(X_ranking_train, y_train)

feature_importance_scores = {feature: 0 for feature in features_to_process}
for feature_name, importance in zip(X_ranking_train.columns, ranking_model.feature_importances_):
    if 'lag_0' in feature_name:
        base_feature = '_'.join(feature_name.split('_')[1:-2])
    else:
        base_feature = '_'.join(feature_name.split('_')[1:-2])
    if base_feature in feature_importance_scores:
        feature_importance_scores[base_feature] += importance

top_features_df = pd.DataFrame(list(feature_importance_scores.items()), columns=['feature', 'importance_score']).sort_values('importance_score', ascending=False)
top_n_features = top_features_df['feature'].head(N_TOP_FEATURES).tolist()
if final_label not in top_n_features:
    top_n_features.append(final_label)

# Modelo Final
final_features_cols = [col for col in X_train.columns if any(col.startswith(tnf) for tnf in top_n_features)]
X_train_final = X_train[final_features_cols]
X_test_final = X_test[final_features_cols]

final_model = RandomForestRegressor(n_estimators=500, random_state=42, n_jobs=-1, max_features=0.33, min_samples_leaf=5)
final_model.fit(X_train_final, y_train)
y_pred = final_model.predict(X_test_final)

```

Figura XXII - Estrutura do Random Forest Regressor usado para testar a previsão do preço do EU ETS

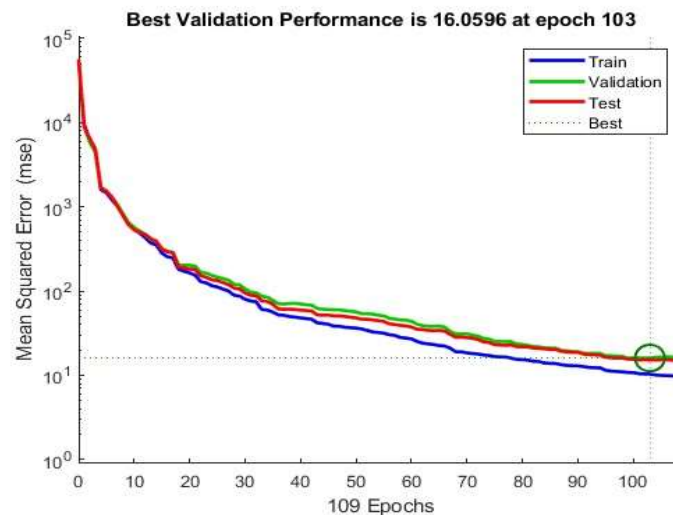
```

all_results_dfs = []
start_date = master_df.index.min()
end_date = master_df.index.max()
current_date = start_date + relativedelta(months=TRAINING_WINDOW_MONTHS)
while current_date < end_date:
    train_start = current_date - relativedelta(months=TRAINING_WINDOW_MONTHS)
    train_end = current_date - pd.Timedelta(days=1)
    test_start = current_date
    test_end = current_date + relativedelta(months=STEP_MONTHS) - pd.Timedelta(days=1)
    if test_end > end_date: test_end = end_date
    train_df = master_df.loc[train_start:train_end]
    test_df = master_df.loc[test_start:test_end]
    if train_df.empty or test_df.empty:
        current_date += relativedelta(months=STEP_MONTHS)
        continue
    y_train = train_df[final_label]
    X_train = train_df.drop(columns=features_to_process)
    y_test = test_df[final_label]
    X_test = test_df.drop(columns=features_to_process)
    scaler = StandardScaler()
    X_train = pd.DataFrame(scaler.fit_transform(X_train), index=X_train.index, columns=X_train.columns)
    X_test = pd.DataFrame(scaler.transform(X_test), index=X_test.index, columns=X_test.columns)
    ranking_features_cols = [col for col in X_train.columns if 'lag_0' in col or (col.split('_')[1].isdigit() and int(col.split('_')[1]) <= N_RANKING_LAGS)]
    ranking_model = RandomForestRegressor(n_estimators=100, random_state=42, n_jobs=-1)
    ranking_model.fit(X_train[ranking_features_cols], y_train)
    feature_importance_scores = {feature: 0 for feature in features_to_process if feature != final_label}
    for feature_name, importance in zip(ranking_features_cols, ranking_model.feature_importances_):
        base_feature = '_'.join(feature_name.split('_')[1:-2])
        if base_feature in feature_importance_scores:
            feature_importance_scores[base_feature] += importance
    top_features_df = pd.DataFrame(list(feature_importance_scores.items()), columns=['feature', 'importance_score']).sort_values('importance_score', ascending=False)
    top_n_features = top_features_df['feature'].head(N_TOP_FEATURES).tolist()
    if final_label not in top_n_features: top_n_features.append(final_label)
    final_features_cols = [col for col in X_train.columns if any(col.startswith(tnf) for tnf in top_n_features)]
    final_model = RandomForestRegressor(n_estimators=500, random_state=42, n_jobs=-1, max_features=0.33, min_samples_leaf=5)
    final_model.fit(X_train[final_features_cols], y_train)
    predictions = final_model.predict(X_test[final_features_cols])
    loop_results_df = pd.DataFrame({'y_true_logdiff': y_test, 'y_pred_logdiff': predictions})
    all_results_dfs.append(loop_results_df)
    current_date += relativedelta(months=STEP_MONTHS)

```

Figura XXIII - Estrutura do Random Forest Regressor com Janela Deslizante para a previsão do EU ETS

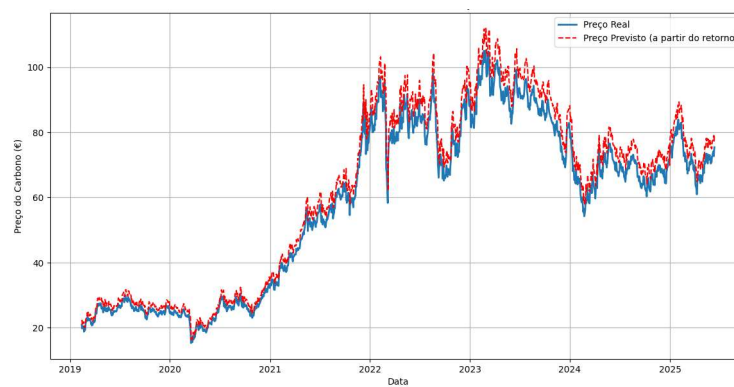
Anexo IX – Função Loss ao longo das várias *epochs* para o treino, validação e teste no Matlab usando uma rede neuronal simples com 100 camadas.



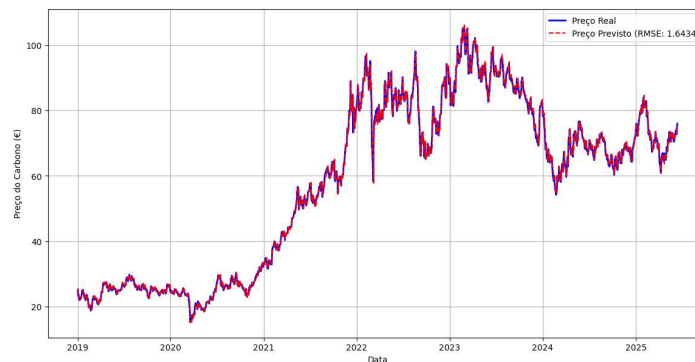
Anexo X – Função *Loss* ao longo das várias *epochs* para o treino, validação e teste no *Jupyter Notebook*, biblioteca *tensorflow* do *python* usando uma rede neuronal do tipo CNN-LSTM seq2seq



Anexo XI – Gráfico que mostra o resultado do preço real vs a previsão do preço usando a CNN-LSTM



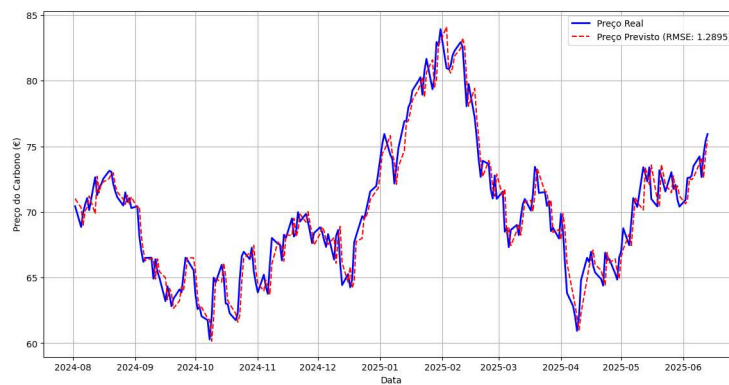
Anexo XII – Gráfico que mostra o resultado do preço real vs a previsão do preço usando o *Random Forest Regressor* com amostra completa



Anexo XIII – Gráfico que mostra o resultado do preço real vs a previsão do preço usando o *Random Forest Regressor* com dados da Fase 3



Anexo XIV – Gráfico que mostra o resultado do preço real vs a previsão do preço usando o *Random Forest Regressor* com dados da Fase 4



Anexo XV – Gráfico que mostra o resultado do preço real vs a previsão do preço usando o a Janela Deslizante no *Random Forest Regressor*

