

UNIVERSIDADE FEDERAL DO PARANÁ

BEATRIZ LEANDRO BONAFINI

MEMORIAL DE PROJETOS: DEFINIÇÕES E QUESTÕES ÉTICAS DA INTELIGÊNCIA
ARTIFICIAL

CURITIBA-PR

2025

BEATRIZ LEANDRO BONAFINI

MEMORIAL DE PROJETOS: DEFINIÇÕES E QUESTÕES ÉTICAS DA INTELIGÊNCIA
ARTIFICIAL

Memorial de Projetos apresentado ao curso de Especialização em Inteligência Artificial Aplicada, Setor de Educação Profissional e Tecnológica, Universidade Federal do Paraná, como requisito parcial à obtenção do título de Especialista em Inteligência Artificial Aplicada.

Área de concentração: *Inteligência Artificial Aplicada*.

Orientador: Prof. Dr. Razer Anthom Nizer Rojas Montañó.

CURITIBA-PR

2025

TERMO DE APROVAÇÃO

Os membros da Banca Examinadora designada pelo Colegiado do Programa de Pós-Graduação Inteligência Artificial Aplicada da Universidade Federal do Paraná foram convocados para realizar a arguição da Monografia de Especialização de **BEATRIZ LEANDRO BONAFINI**, intitulada: **MEMORIAL DE PROJETOS: DEFINIÇÕES E QUESTÕES ÉTICAS DA INTELIGÊNCIA ARTIFICIAL**, que após terem inquirido a aluna e realizada a avaliação do trabalho, são de parecer pela sua aprovação no rito de defesa.

A outorga do título de especialista está sujeita à homologação pelo colegiado, ao atendimento de todas as indicações e correções solicitadas pela banca e ao pleno atendimento das demandas regimentais do Programa de Pós-Graduação.

Curitiba, 22 de Agosto de 2025.



RAZER ANTHOM NIZER ROJAS MONTAÑO
Presidente da Banca Examinadora



RAFAELA MANTOVANI FONTANA
Avaliador Interno (UNIVERSIDADE FEDERAL DO PARANÁ)

AGRADECIMENTOS

A Deus, por conceder-me forças para perseverar, por não me desamparar nos momentos de dificuldade e, igualmente, por permitir que eu reconhecesse, nos instantes de alegria, a perfeição de Seus planos para a minha vida. Durante este percurso, encontrei conforto na Palavra, razão pela qual dedico este trabalho, primeiramente, a Ele.

Aos meus pais, Mari e Lúcio, exemplos de retidão, honestidade e determinação, que me orientaram com amor e dedicação ao longo desta jornada. Sem o apoio incondicional deles, não teria sido possível alcançar esta etapa.

Ao meu irmão Guilherme, cuja influência foi determinante para a minha aproximação da área da computação, e que, ainda na infância, me auxiliou no aprendizado dos primeiros algoritmos.

Ao meu noivo, Eric, pela presença constante, por me motivar, de forma contínua, a perseguir meus objetivos.

Aos colegas de curso, pela disponibilidade em colaborar com as disciplinas, trabalhos e demais desafios que compuseram esta especialização.

À instituição Centro de Estudos e Sistemas Avançados do Recife (CESAR), por oportunizar meu contínuo aprimoramento profissional.

Aos docentes da Especialização em Inteligência Artificial Aplicada, pelo conhecimento transmitido ao longo do curso. De modo especial, ao professor Razer Anthom, pela dedicação, pela disponibilidade em auxiliar os discentes e pela orientação prestada nesta etapa conclusiva.

"Senhor, fazei-me instrumento de vossa paz. Onde houver ódio, que eu leve o amor; Onde houver ofensa, que eu leve o perdão; Onde houver discórdia, que eu leve a união; Onde houver dúvida, que eu leve a fé; Onde houver erro, que eu leve a verdade; Onde houver desespero, que eu leve a esperança; Onde houver tristeza, que eu leve a alegria; Onde houver trevas, que eu leve a luz.

Ó Mestre, fazei que eu procure mais: Consolar, que ser consolado; Compreender, que ser compreendido; Amar, que ser amado.

Pois é dando que se recebe; É perdando que se é perdoado; E é morrendo que se vive para a vida eterna. Amém"

Oração de São Francisco de Assis

RESUMO

O conceito de Inteligência Artificial (IA) é amplamente difundido, mas permanece subjetivo devido às múltiplas definições e abordagens adotadas por diferentes autores. A partir da disciplina de Introdução à IA, foram analisadas quatro grandes perspectivas: pensar como humanos, agir como humanos, pensar racionalmente e agir racionalmente. Essas abordagens refletem tanto a tentativa de replicar o raciocínio humano quanto a busca por desempenho racional e eficiente em máquinas. Além das definições técnicas, foram discutidas questões éticas e filosóficas relacionadas ao avanço da IA, como a substituição de mão de obra, o uso indevido da tecnologia, a responsabilização por erros e os riscos à segurança e aos direitos civis. Também se destacou a importância da governança e da definição de salvaguardas para garantir o uso responsável da IA. Por fim, enfatiza-se a necessidade de uma abordagem interdisciplinar que una tecnologia, ética, direito e ciências sociais para mitigar os riscos e promover uma IA transparente, segura e alinhada aos valores humanos.

Palavras-chave: Inteligência Artificial; Ética em IA; Agentes Inteligentes; Responsabilidade Algorítmica; Impactos Sociais da IA; Definições de IA.

ABSTRACT

The concept of Artificial Intelligence (AI) is widely disseminated but remains subjective due to the multiple definitions and approaches adopted by different authors. Based on discussions from the Introduction to AI course, four main perspectives were analyzed: thinking like humans, acting like humans, thinking rationally, and acting rationally. These approaches reflect both the attempt to replicate human reasoning and the pursuit of rational and efficient machine performance. In addition to technical definitions, ethical and philosophical issues related to the advancement of AI were addressed, such as the replacement of human labor, the misuse of technology, accountability for errors, and risks to security and civil rights. The importance of governance and the implementation of safeguards was emphasized to ensure the responsible use of AI. Finally, the need for an interdisciplinary approach—integrating technology, ethics, law, and social sciences—is highlighted as essential to mitigate risks and promote transparent, safe, and human-aligned AI systems.

Keywords: Artificial Intelligence; Ethics in AI; Intelligent Agents; Algorithmic Responsibility; Social Impacts of AI; Definitions of AI.

LISTA DE FIGURAS

1.1	Exercício 2: Busca Heurística - Melhor Caminho para Bucharest	18
1.2	Exercício 4: Rede Neural Artificial.	20
1.3	Resolução Exercício 2: Caminho de Lugoj até Bucharest	22
2.1	Resolução Exercício 2: Distribuição da Quantidade de Carros por Marca.	27
2.2	Resolução Exercício 2: Distribuição da Quantidade de Carros por Tipo de Engrenagem.	28
2.3	Resolução Exercício 2: Evolução da Média de Preço dos Carros ao Longo de 2022	29
2.4	Resolução Exercício 2: Média de Preços dos Carros por Marca e Tipo de Engrenagem Agrupados por Tipo.	30
2.5	Resolução Exercício 2: Média de Preços dos Carros por Marca e Tipo de Engrenagem Agrupados por Marca.	31
2.6	Resolução Exercício 2: Média de Preços dos Carros por Marca e Tipo de Combustível Agrupados por Marca.	32
2.7	Resolução Exercício 2: Média de Preços dos Carros por Marca e Tipo de Combustível Agrupados por Tipo.	33
2.8	Resolução Exercício 3: Boxplot da Distribuição dos Preços Médios	34
2.9	Resolução Exercício 3: Mapa de Correlação entre as Variáveis Numéricas	35
7.1	Admissão - Gráfico de resíduos (melhor modelo)	85
7.2	Biomassa - Gráfico de resíduos (melhor modelo)	85
8.1	Classificação de Imagens com CNN: Gráfico de Perda.	99
8.2	Classificação de Imagens com CNN: Gráfico de Acurácia	100
8.3	Classificação de Imagens com CNN: Matriz de Confusão	100
8.4	Classificação de Imagens com CNN: Resultado da Predição.	101
8.5	Detector de SPAM com RNN: Gráfico de Perda	103
8.6	Detector de SPAM com RNN: Gráfico de Acurácia	103
8.7	Gerando dados com GAN: Geração de Imagem Ruidosa	105
8.8	Gerando dados com GAN: Checkpoint da Geração de Dígitos.	108
8.9	Tranformer: Codificação Posicional	111
9.1	Visão geral das Ferramentas Utilizadas pelo AirBnb.	122
13.1	Exercício 1: Gráfico de Perda.	156
13.2	Exercício 1: Gráfico de Acurácia	156
13.3	Exercício 1: Matriz de Confusão	157
13.4	Exercício 1: Apresentando Classificação Incorreta.	157
13.5	Exercício 2: Gráfico de Perda.	159
13.6	Exercício 2: Gráfico de RMSE	160

13.7	Exercício 2: Gráfico de R2	160
13.8	Exercício 3: Gráfico de Perda.	162
13.9	Exercício 3: Imagem Baixada da URL Fornecida	163
13.10	Exercício 3: Resultado Obtido no Processo de Deep Dream	166
13.11	Exercício 3: Resultado Obtido no Processo de Deep Dream à Oitava	166
14.1	Evolução da Produção do Conteúdo na Netflix	169
15.1	Exercício 1: Rota Inicial Aleatória, Distância = 5159,51.	176
15.2	Exercício 1: Rota Após Evolução, Distância = 1154,43	177
15.3	Exercício 2: Projeção PCA 3D	179
15.4	Exercício 2: Projeção PCA 2D	179
15.5	Exercício 2: Dendograma - Distâncias Originais.	180
15.6	Exercício 2: Dendograma - Distâncias PCA 2D	181

LISTA DE TABELAS

1.1	Valores de h_{DLR} — distâncias em linha reta para Bucareste..	19
2.1	Exercício 2 de LPA: Metados da Base de Dados Carros Brasil	23
5.1	Desempenho dos modelos de regressão nas etapas de treino e teste..	55
7.1	Veículos - Resultados por técnica com parâmetros, acurácia e matrizes de confusão.	82
7.2	Diabetes - Resultados por técnica com parâmetros, acurácia e matriz de confusão.	83
7.3	Admissão - Resultados por técnica com parâmetros e métricas (R^2 , S_{yx} , coeficiente de Pearson, RMSE e MAE).	84
7.4	Amostra com GRE, TOEFL, rating universitário e métricas preditivas.	84
7.5	Biomassa - Resultados por técnica com parâmetros e métricas (R^2 , S_{yx} , coeficiente de Pearson, RMSE e MAE).	85
7.6	Amostra com variáveis dap , h , Me e predição da RNA.	85
7.7	K-means (médias por cluster) - Parte 1	86
7.8	K-means (médias por cluster) - Parte 2	86
7.9	Regras de associação — top 10 linhas.	86

SUMÁRIO

1	INTELIGÊNCIA ARTIFICIAL	11
1.1	DEFINIÇÃO	11
1.2	QUESTÕES ÉTICAS E FILOSÓFICAS DA IA	13
	REFERÊNCIAS	17
	APÊNDICE 1 – INTRODUÇÃO À INTELIGÊNCIA ARTIFICIAL	18
	APÊNDICE 2 – LINGUAGEM DE PROGRAMAÇÃO APLICADA	23
	APÊNDICE 3 – LINGUAGEM R.	40
	APÊNDICE 4 – ESTATÍSTICA APLICADA I	48
	APÊNDICE 5 – ESTATÍSTICA APLICADA II.	54
	APÊNDICE 6 – ARQUITETURA DE DADOS	72
	APÊNDICE 7 – APRENDIZADO DE MÁQUINA	82
	APÊNDICE 8 – DEEP LEARNING	98
	APÊNDICE 9 – BIG DATA	121
	APÊNDICE 10 – VISÃO COMPUTACIONAL	124
	APÊNDICE 11 – ASPECTOS FILOSÓFICOS E ÉTICOS DA IA	146
	APÊNDICE 12 – GESTÃO DE PROJETOS DE IA.	151
	APÊNDICE 13 – FRAMEWORKS DE INTELIGÊNCIA ARTIFICIAL. .	153
	APÊNDICE 14 – VISUALIZAÇÃO DE DADOS E STORYTELLING . . .	167
	APÊNDICE 15 – TÓPICOS EM INTELIGÊNCIA ARTIFICIAL	170

1 INTELIGÊNCIA ARTIFICIAL

1.1 DEFINIÇÃO

Hoje em dia, o termo *Inteligência Artificial* (IA) é muito comum. No entanto, o conceito é subjetivo, pois é usado e aplicado em diferentes contextos, levando a várias interpretações e definições. Durante a disciplina de Introdução à IA, ministrada por Montañó (2024), foi possível observar que diferentes autores oferecem definições conflitantes em relação aos processos de pensamento e raciocínio, bem como ao comportamento exibido por máquinas. Ao analisar o processo de replicação do pensamento e raciocínio humanos por máquinas, alguns autores, definem IA como uma máquina que pensa como um ser humano ou pensa racionalmente. Por outro lado, ao observar o comportamento das máquinas, outros autores definem IA como aquela capaz de agir como um ser humano ou agir racionalmente. Antes de aprofundar cada definição, é importante reconhecer que os pensamentos e ações humanas nem sempre são racionais. Segundo Norvig e Russell (2022), as definições de IA podem ser agrupadas em quatro grandes abordagens: pensar como humanos, agir como humanos, pensar racionalmente e agir racionalmente. Nos parágrafos seguintes, detalha-se cada uma dessas abordagens.

Pensando como humanos: Haugeland (1985) argumenta que a IA está associada ao pensamento humano, visto que se trata de um esforço novo e intrigante para fazer os computadores pensarem como máquinas com mentes, no sentido total e literal. Por outro lado, Bellman (1978) define IA como atividades que associamos ao pensamento humano, como tomada de decisão, resolução de problemas e aprendizado. De acordo com essa perspectiva, o objetivo principal é implementar fielmente o processo de pensamento em um computador por meio de programas. No entanto, é importante lembrar que, embora algoritmos possam produzir resultados corretos, isso não significa que eles simulem o comportamento humano. Para determinar se um programa pensa, devemos primeiro entender como os seres humanos pensam.

Agindo como humanos: A IA, segundo Kurzweil (1990), pode ser definida como “a arte de criar máquinas que executam funções que exigem inteligência quando executadas por pessoas”. De forma similar, Rich e Knight (1991) a definem como “o estudo de como os computadores podem fazer tarefas que hoje são melhor desempenhadas pelas pessoas”. Nesse contexto, o objetivo da IA seria imitar o comportamento humano, sem necessariamente implicar

a implementação de raciocínio. A relevância reside no fato de o resultado fornecido ser similar ao produzido por um ser humano. Turing (1950) propõe uma forma de medir o nível satisfatório de inteligência por meio do Teste de Turing. O teste consiste em uma conversa entre um humano, outro humano e uma máquina, utilizando linguagem natural. Todos os participantes estão separados. Se não for possível distinguir com segurança a máquina do humano, a máquina é considerada aprovada no teste. É importante ressaltar que, embora a máquina forneça respostas, o teste não verifica a correção dessas respostas, mas sim a proximidade com as respostas humanas.

Pensando racionalmente: De acordo com Charniak e McDermott (1985), a IA consiste no “estudo das faculdades mentais por meio de modelos computacionais”. Winston (1992), por sua vez, define a IA como “o estudo das computações que tornam possível perceber, raciocinar e agir”. O objetivo, nesse contexto, é a modelagem do pensamento correto, ou seja, a modelagem do processo de raciocínio, a fim de fornecer um resultado logicamente válido. Entretanto, é fundamental compreender como se define um pensamento correto. Assim Aristóteles apresenta um conjunto de frases que formam silogismos. Os silogismos possibilitam a construção do pensamento correto, sendo compostos por três sentenças: duas afirmativas e uma conclusiva. Um exemplo seria: Sócrates é um homem (afirmativa 1); todos os homens são mortais (afirmativa 2); portanto, Sócrates é mortal.

Agindo Racionalmente: para Poole, Mackworth e Goebel (1998), a “Inteligência Computacional é o estudo do projeto de agentes inteligentes”. Nilsson (1998) diz que a IA “está relacionada a um desempenho inteligente de artefatos”. O objetivo, nesse contexto, é implementar agentes capazes de responder a situações e buscar o melhor resultado possível. Um agente, em sua essência, age sob controle autônomo, percebe seu ambiente, adapta-se às mudanças e é capaz de criar e perseguir metas. Durante a disciplina, foi demonstrado que este é o conceito ideal de IA apresentado por Norvig e Russell (2022), em que “a IA se concentra no estudo e na construção de agentes que fazem a coisa certa”.

Diante das diversas definições apresentadas sobre IA, é compreensível que exista certa confusão em relação ao seu real significado e aplicação. A multiplicidade de perspectivas, que vai desde o desempenho inteligente de artefatos até a construção de agentes que fazem a coisa certa, contribui para essa complexidade. Atualmente, observa-se uma tendência crescente de associar a IA a praticamente a maior parte das soluções tecnológicas, criando a impressão de que qualquer inovação deva obrigatoriamente envolver essa tecnologia. No entanto, é importante destacar que nem todas as soluções necessitam da aplicação da IA. Muitas vezes, abordagens

mais simples e tradicionais podem ser mais eficientes, econômicas e adequadas ao contexto. O uso da IA deve ser considerado quando realmente agrega valor, resolve problemas complexos ou otimiza processos de forma significativa. Compreender essa distinção é fundamental para evitar o uso incorreto do termo e para garantir que a tecnologia seja utilizada de maneira ética, responsável e eficaz.

1.2 QUESTÕES ÉTICAS E FILOSÓFICAS DA IA

Ao longo da disciplina de Montañó (2024), foram explorados diversos pontos relacionados às preocupações sobre o avanço da IA na sociedade contemporânea. Os debates concentraram-se em temas como a substituição da mão de obra humana por IA, o uso indevido da tecnologia, questões de responsabilização, riscos de destruição em massa e desafios relacionados ao plágio.

No que diz respeito à substituição da mão de obra, Montañó (2024) destaca que discussões semelhantes já ocorreram em momentos históricos marcados pelo surgimento de novas tecnologias, como a calculadora, a internet e a uberização dos serviços. Ressalta-se que a IA é capaz de executar tarefas que, até recentemente, não existiam ou eram extremamente onerosas para os seres humanos. Entre essas tarefas, destacam-se a análise financeira complexa, a detecção de fraudes, a realização de atividades repetitivas, a manipulação de materiais perigosos e até a exploração de outros planetas. Considerando o avanço exponencial da internet e a imensa geração de dados, torna-se evidente a necessidade de agentes capazes de processar e extrair informações em uma escala que ultrapassa as capacidades humanas. A IA, nesse contexto, surge como uma ferramenta indispensável para lidar com desafios que envolvem grandes volumes de dados e complexidade analítica. Outro ponto neste contexto é que há o surgimento de outras posições de trabalho, melhor remuneradas e qualificadas. Um exemplo disto é o surgimento de funções como Arquiteto de IA, Cientista de Dados e Engenheiro de *Machine Learning*.

No que diz respeito ao uso indevido da IA, é fundamental destacar que essa preocupação é recorrente em diversas áreas da ciência. Um exemplo discutido durante a disciplina refere-se à tecnologia nuclear, que pode ser aplicada tanto para fins benéficos, como a geração de energia, quanto para propósitos indesejáveis, como a fabricação de armas nucleares. Além disso, o uso da IA em contextos militares também é um tema que gera intensos debates, evidenciando a necessidade de uma reflexão ética sobre suas aplicações.

Durante o período da Especialização em IA Aplicada, a aluna Beatriz Leandro Bonafini foi selecionada para representar o Brasil em uma importante discussão internacional sobre o uso responsável da IA, realizada em junho de 2024, em Lisboa, Portugal. O evento foi promovido pela UNODA (*United Nations Office for Disarmament Affairs*) em parceria com o SIPRI (*Stockholm International Peace Research Institute*) e contou com a participação de novos pesquisadores de 15 países, todos atuantes na área de IA. O principal objetivo do encontro foi fomentar o debate sobre o futuro da IA no contexto da segurança e soberania das nações. Durante as discussões, temas que foram abordados na disciplina da especialização foram vistos, como o controle de Veículos Aéreos Não Tripulados (VANTs) e as estratégias que órgãos governamentais podem adotar para enfrentar desafios relacionados à responsabilização pelo uso da IA. Destacou-se, ainda, a importância da implementação de salvaguardas em soluções de IA, que, conforme também discutido no curso, deixaram de ser opcionais para se tornarem medidas obrigatórias. Em relação à responsabilização por danos causados por sistemas de IA, enfatizou-se que o poder de decisão deve estar nas mãos daqueles que determinam a finalidade do uso da tecnologia, garantindo, assim, maior segurança e ética nas suas aplicações.

Além do uso da IA para fins militares, existem ainda outras formas em que a IA é utilizada de maneira indesejada, tais como:

- Disseminação de *fake news*: A IA pode automatizar a disseminação de notícias falsas, ampliando rapidamente o alcance da desinformação e dificultando a verificação dos fatos;
- *Deepfake* pornográfico: O uso de IA para criar vídeos falsificados de conteúdo adulto representa uma séria violação da privacidade, prejudicando a reputação das vítimas e levantando questões éticas e legais;
- Reforço de preconceitos: Algoritmos de IA podem intensificar preconceitos existentes, caso sejam treinados com dados enviesados, contribuindo para discriminação em áreas como recrutamento, crédito e segurança pública; Plágio acadêmico: Ferramentas de IA podem ser usadas para gerar textos acadêmicos sem a devida originalidade, facilitando casos de plágio e comprometendo a integridade da pesquisa científica;
- Criação de discussões fictícias: *Bots* alimentados por IA podem simular interações humanas em redes sociais, criando debates artificiais que manipulam a percepção da opinião pública;

- Levantamento de *hashtags* em redes sociais: Algoritmos podem ser empregados para impulsionar *hashtags* de forma artificial, influenciando tendências e direcionando a atenção do público para determinados assuntos;
- Manipulação da opinião pública: A IA pode segmentar e direcionar conteúdos com o objetivo de influenciar comportamentos e decisões, especialmente em contextos eleitorais e políticos;
- Direcionamento de embates políticos: O uso de IA para polarizar o debate público, promovendo narrativas extremistas, contribuindo para o aumento da tensão em sociedades democráticas;
- Perda dos direitos civis: com a facilitação da vigilância em massa é possível que pessoas tenham perda da privacidade dado histórico de uso (análise de e-mails, ligações e transações).

De quem é a responsabilidade quando a IA comete um erro? Durante a disciplina, foi abordado um exemplo interessante: um médico, auxiliado por uma IA, fornece um diagnóstico incorreto, gerando o questionamento sobre a quem cabe essa responsabilidade. Nesse caso, foi destacado que o médico mantém o controle da decisão. Enfatizou-se que os sistemas especialistas são ferramentas de apoio, não devendo ser responsáveis pelas decisões finais. O erro é inerente à IA, e isso deve ser considerado sempre que um resultado impreciso for fornecido por ela. O conteúdo apresentado também destacou que, embora sistemas possam superar o desempenho humano em algumas áreas, se um profissional optar por não seguir as recomendações do sistema, ele será o responsável pelas consequências. Deve-se lembrar que robôs e softwares operados por IA são produtos.

Será que uma IA tem a capacidade de aniquilar a raça humana? Esse foi outro tópico abordado ao longo da disciplina, evidenciando que qualquer tecnologia poderia, em teoria, destruir a humanidade. De forma hipotética, imagina-se que a IA poderia se tornar consciente e tomar a decisão equivocada de que os seres humanos geram sofrimento. Assim, para eliminar o sofrimento humano, ela poderia optar por destruir a própria humanidade. Nesse contexto, é necessária a definição de uma função de utilidade, estimativas de estados incorretos e uma função de aprendizagem. A função de utilidade tem o objetivo de verificar se as ações executadas pela IA correspondem à sua real utilidade em determinado contexto. Já a estimativa do estado da

IA busca identificar erros cometidos e implementar ponderações adequadas quando esses erros ocorrem. A função de aprendizagem consiste na validação do que ela está aprendendo.

Por fim, é importante ressaltar que a forma como as tecnologias vêm evoluindo dentro da IA torna-se um desafio para os profissionais da área. Os desafios não se limitam apenas à aquisição de dados, ao desenvolvimento e ao treinamento de uma IA, mas também abrangem os impactos que essas soluções podem gerar em diversas esferas da sociedade. Além disso, destaca-se a necessidade de uma abordagem ética e responsável no desenvolvimento dessas tecnologias, considerando questões como privacidade, segurança e o potencial de viés nos algoritmos. A governança da IA torna-se fundamental para garantir que seu uso seja benéfico e alinhado aos valores humanos, mitigando riscos e promovendo a transparência nas decisões automatizadas. Por conseguinte, a interdisciplinaridade entre tecnologia, filosofia, direito e ciências sociais é essencial para compreender e gerenciar os efeitos da IA no cotidiano, assegurando que o progresso tecnológico contribua para o bem-estar coletivo.

REFERÊNCIAS

- BELLMAN, R. E. *An Introduction to Artificial Intelligence: Can Computers Think?* [S.l.]: Boyd & Fraser Publishing Company, 1978.
- CHARNIAK, E.; MCDERMOTT, D. *Introduction to Artificial Intelligence*. [S.l.]: Addison-Wesley, 1985.
- HAUGELAND, J. *Artificial intelligence: The very idea*. [S.l.]: MIT Press, 1985.
- KURZWEIL, R. *The Age of Intelligent Machines*. [S.l.]: MIT Press, 1990.
- MONTAÑO, R. A. N. R. Introdução ao aprendizado de máquina. Aula ministrada no Curso de Especialização em Inteligência Artificial Aplicada, Universidade Federal do Paraná. Não há registro público. 2024.
- NILSSON, N. J. *Artificial Intelligence: A New Synthesis*. [S.l.]: Morgan Kaufmann, 1998.
- NORVIG, P.; RUSSELL, S. *Artificial Intelligence: A Modern Approach, 4th US ed.* [S.l.]: Pearson. <http://aima.cs.berkeley.edu>, 2022.
- POOLE, D.; MACKWORTH, A. K.; GOEBEL, R. *Computational Intelligence: A Logical Approach*. New York: Oxford University Press, 1998.
- RICH, E.; KNIGHT, K. *Artificial Intelligence*. 2. ed. [S.l.]: McGraw-Hill, 1991.
- TURING, A. M. Computing machinery and intelligence. *Mind*, v. 59, n. 236, p. 433–460, 1950.
- WINSTON, P. H. *Artificial Intelligence*. 3. ed. [S.l.]: Addison-Wesley, 1992.

APÊNDICE 1 – INTRODUÇÃO À INTELIGÊNCIA ARTIFICIAL

A - ENUNCIADO 1 ChatGPT

- (6,25 pontos) Pergunte ao ChatGPT o que é Inteligência Artificial e cole aqui o resultado.
- (6,25 pontos) Dada essa resposta do ChatGPT, classifique usando as 4 abordagens vistas em sala. Explique o porquê.
- (6,25 pontos) Pesquise sobre o funcionamento do ChatGPT (sem perguntar ao próprio ChatGPT) e escreva um texto contendo no máximo 5 parágrafos. Cite as referências.
- (6,25 pontos) Entendendo o que é o ChatGPT, classifique o próprio ChatGPT usando as 4 abordagens vistas em sala. Explique o porquê.

2 Busca Heurística

Realize uma busca utilizando o algoritmo A* para encontrar o melhor caminho para chegar a **Bucharest** partindo de **Lugoj**. Construa a árvore de busca criada pela execução do algoritmo apresentando os valores de $f(n)$, $g(n)$ e $h(n)$ para cada nó. Utilize a heurística de distância em linha reta, que pode ser observada na tabela abaixo.

Essa tarefa pode ser feita em uma **ferramenta de desenho**, ou até mesmo no **papel**, desde que seja digitalizada (foto) e convertida para PDF.

- (25 pontos) Apresente a árvore final, contendo os valores, da mesma forma que foi apresentado na disciplina e nas práticas. Use o formato de árvore, não será permitido um formato em blocos, planilha, ou qualquer outra representação.

NÃO É NECESSÁRIO IMPLEMENTAR O ALGORITMO.

Figura 1.1 – EXERCÍCIO 2: BUSCA HEURÍSTICA - MELHOR CAMINHO PARA BUCHAREST

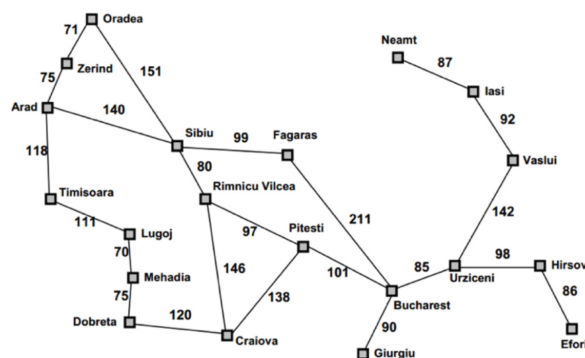


Tabela 1.1: Valores de h_{DLR} — distâncias em linha reta para Bucareste.

Arad	366	Mehadia	241
Bucareste	0	Neamt	234
Craiova	160	Oradea	380
Drobeta	242	Pitesti	100
Eforie	161	Rimnicu Vilcea	193
Fagaras	176	Sibiu	253
Giurgiu	77	Timisoara	329
Hirsova	151	Urziceni	80
Iasi	226	Vaslui	199
Lugoj	244	Zerind	374

3 Lógica

Verificar se o argumento lógico é válido:

Se as uvas caem, então a raposa as come

Se a raposa as come, então estão maduras

As uvas estão verdes ou caem

Logo

A raposa come as uvas se e somente se as uvas caem

Deve ser apresentada uma prova, no mesmo formato mostrado nos conteúdos de aula e nas práticas.

Dicas:

1. Transformar as afirmações para lógica:

p : as uvas caem

q : a raposa come as uvas

r : as uvas estão maduras

2. Transformar as três primeiras sentenças para formar a base de conhecimento:

$R1 : p \rightarrow q$

$R2 : q \rightarrow r$

$R3 : \neg r \vee p$

3. Aplicar equivalências e regras de inferência para se obter o resultado esperado. Isto é, com essas três primeiras sentenças devemos derivar $q \leftrightarrow p$. Cuidado com a ordem em que as fórmulas são geradas.

Equivalência Implicação: $(\alpha \rightarrow \beta)$ equivale a $(\neg \alpha \vee \beta)$

Silogismo Hipotético: $\alpha \rightarrow \beta, \beta \rightarrow \gamma \vdash \alpha \rightarrow \gamma$

Conjunção: $\alpha, \beta \vdash \gamma \wedge \beta$

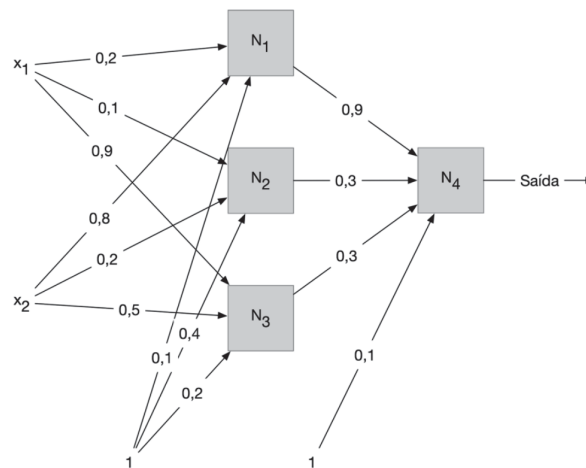
Equivalência Bicondicional: $(\alpha \leftrightarrow \beta)$ equivale a $(\alpha \rightarrow \beta) \wedge (\beta \rightarrow \alpha)$

- a) **(25 pontos)** Deve-se mostrar todos os passos e regras aplicadas, no mesmo formato apresentado nas aulas e nas práticas. As equivalências e regras necessárias estão descritas acima e no material.

4 Redes Neurais Artificiais

Seja a RNA da figura abaixo:

Figura 1.2 – EXERCÍCIO 4: REDE NEURAL ARTIFICIAL



Os neurônios N_1, N_2 e N_3 possuem função de ativação linear. Já N_4 possui função de ativação tangente hiperbólica (pesquise a fórmula e aplique). Dada a entrada $x_1 = -3, x_2 = 1$, dê:

- (6,25 pontos)** Valor de saída do neurônio N_1
- (6,25 pontos)** Valor de saída do neurônio N_2
- (6,25 pontos)** Valor de saída do neurônio N_3
- (6,25 pontos)** Valor de saída da rede como um todo

B - RESOLUÇÃO

1 ChatGPT

- a) Inteligência Artificial (IA) refere-se à capacidade de uma máquina ou programa de computador realizar tarefas que normalmente exigiria inteligência humana. Essas tarefas incluem aprendizado, raciocínio, percepção, reconhecimento de padrões, compreensão da linguagem natural e tomada de decisões. O objetivo da Inteligência Artificial é criar sistemas que possam executar funções complexas de forma autônoma, imitando ou replicando a inteligência humana. Existem dois tipos principais de IA: IA fraca (ou estreita) e IA forte (ou geral).

1. IA fraca (estreita): Refere-se a sistemas especializados em realizar tarefas específicas sem possuir a capacidade de generalizar para diferentes domínios. Exemplos incluem assistentes virtuais, reconhecimento de voz, carros autônomos e sistemas de recomendação.

2. IA forte (geral): Representa um nível mais avançado de inteligência, comparável à inteligência humana em termos de habilidades cognitivas e adaptação a diferentes contextos. Até o momento, a IA forte permanece mais uma aspiração do que uma realidade prática.

As técnicas comuns em IA incluem aprendizado de máquina (*machine learning*), redes neurais, processamento de linguagem natural, visão computacional e algoritmos de otimização. O campo da Inteligência Artificial continua a evoluir rapidamente, com aplicações em diversas áreas, como saúde, finanças, transporte, manufatura e entretenimento.

- b) É categorizável na abordagem "agir como seres humanos", pois no texto o chat GPT ressaltou a imitação de habilidades inerentemente humanas ao tentar definir o conceito de Inteligência Artificial.
- c) Como o ChatGPT funciona: ChatGPT Sigla para *Generative Pré-Trained Transformer* (Transformador pré-treinado generativo, em tradução livre) é um sistema ou um modelo de linguagem baseado em inteligência artificial mais especificamente com as técnicas de aprendizado de máquinas, redes neurais, aprendizado por reforço com *feedback* humano (sigla em inglês RHLF) e processamento de linguagem natural usados com o foco em diálogos virtuais. A seguir vamos falar sobre cada uma das partes do algoritmo GPT. O ChatGPT usa como fonte de dados para manter os seus diálogos virtuais e o contexto das mesmas um conjunto de informações disponíveis e acessíveis principalmente através da internet.

O generativo pré-treinado, o GP de GPT, é uma técnica onde ChatGPT recebe uma grande quantidade de regras de linguagens e dados não rotulados. O GPT é então deixado sem supervisão para aprender sozinho sobre as regras e os relacionamentos que governam as linguagens.

A letra T de GPT, *Transformer architecture*, refere-se à arquitetura do modelo, a qual é baseada na arquitetura *Transformer*. Agora, todos os dados pré-treinados são usados para criar uma rede neural de aprendizagem profunda, permitindo ao ChatGPT aprender padrões e relacionamentos de textos, dando ao ChatGPT a habilidade de criar respostas como humanos e de prever qual texto vem em seguida em uma dada sentença. O *Transformer* é capaz de ler cada palavra de uma sentença e comparar cada palavra com as outras palavras da mesma sentença. Isso permite que ele direcione a sua atenção à palavra mais importante da sentença. Esse processo é conhecido como *self-attention*. Com isso, é importante ressaltar que o ChatGPT não entende o que é dito e o que ele responde.

No início, a rede neural do ChatGPT não estava inteiramente segura para ser lançada ao público, pois ela foi treinada e sumariamente através de dados da internet aberta sem nenhuma orientação. Então para que o ChatGPT pudesse dar respostas coerentes, sensatas e seguras ao público, foi otimizado com uma técnica chamada aprendizado por reforço com *feedback* humano (sigla em inglês RHLF). Basicamente são demonstrados dados que guiam a rede neural em como responder adequadamente às solicitações humanas. O RHLF, mesmo não sendo um aprendizado supervisionado puro, permite ao GPT ser efetivamente *fine-tuned*.

Outra técnica igualmente utilizada é o *natural language processing* (NLP), que é o processo de fazer um inteligência artificial a entender as regras e a sintaxe de uma linguagem.

Referências

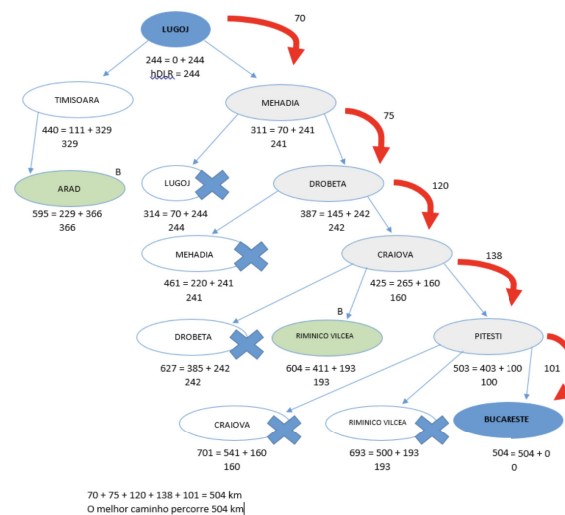
<https://zapier.com/blog/how-does-chatgpt-work/>

<https://investnews.com.br/guias/chatgpt/>

- d) Pensar racionalmente, pois fornece respostas com base em padrões aprendidos a partir de uma vasta gama de dados. Ele não possui emoções, experiências pessoais ou compreensão profunda das nuances humanas. E ele não entende o que é dito e o que ele responde.

2 Busca Heurística

Figura 1.3 – RESOLUÇÃO EXERCÍCIO 2: CAMINHO DE LUGOJ ATÉ BUCHAREST



3 Lógica

$R4 : r \rightarrow p$ - Equivalência da implicação em R3

$R5 : q \rightarrow p$ - Silogismo hipotético em R2 e R4

$R6 : p \rightarrow q \wedge q \rightarrow p$ - Conjunção em R1 e R5

$R7 : q \leftrightarrow p$ - Equivalência bicondicional em R6

4 Redes Neurais Artificiais

a) $N_1 = 0,3$ função linear onde $fa(u) = u - 3 * 0,2 + 1 * 0,8 + 1 * 0,1 = 0,3$

b) $N_2 = 0,3 | -3 * 0,1 + 1 * 0,2 + 1 * 0,4 = 0,3$

c) $N_3 = -2 | -3 * 0,9 + 1 * 0,5 + 1 * 0,2 = -2$

d) $Sada = -0,139$

$$0,3 * 0,9 + 0,3 * 0,3 - 2 * 0,3 + 1 * 0,1 = -0,14$$

$$TangH = (EXP(-0,14) - EXP(-(-0,14)))/(EXP(-0,14) + EXP(-(-0,14))) = -0,139$$

APÊNDICE 2 – LINGUAGEM DE PROGRAMAÇÃO APLICADA

A - ENUNCIADO

Nome da base de dados do exercício: `precos_carros_brasil.csv`

Informações sobre a base de dados:

Dados dos preços médios dos carros brasileiros, das mais diversas marcas, no ano de 2021, de acordo com dados extraídos da tabela FIPE (Fundação Instituto de Pesquisas Econômicas). A base original foi extraída do site Kaggle (Acesse aqui a base original). A mesma foi adaptada para ser utilizada no presente exercício.

Observação: As variáveis `fuel`, `gear` e `engine_size` foram extraídas dos valores da coluna `model`, pois na base de dados original não há coluna dedicada a esses valores. Como alguns valores do modelo não contêm as informações do tamanho do motor, este conjunto de dados não contém todos os dados originais da tabela FIPE.

Metadados:

Tabela 2.1: Exercício 2 de LPA: Metadados da Base de Dados Carros Brasil

Nome do campo	Descrição
<code>year_of_reference</code>	O preço médio corresponde a um mês de ano de referência.
<code>month_of_reference</code>	O preço médio corresponde a um mês de referência, ou seja, a FIPE atualiza sua tabela mensalmente.
<code>fipe_code</code>	Código único da FIPE.
<code>authentication</code>	Código de autenticação único para consulta no site da FIPE.
<code>brand</code>	Marca do carro.
<code>model</code>	Modelo do carro.
<code>fuel</code>	Tipo de combustível do carro.
<code>gear</code>	Tipo de engrenagem do carro.
<code>engine_size</code>	Tamanho do motor em centímetros cúbicos.
<code>year_model</code>	Ano do modelo do carro. Pode não corresponder ao ano de fabricação.
<code>avg_price</code>	Preço médio do carro, em reais.

Atenção: ao fazer o download da base de dados, selecione o formato `.csv`. É o formato que será considerado correto na resolução do exercício.

1 Análise Exploratória dos Dados

A partir da base de dados `precos_carros_brasil.csv`, execute as seguintes tarefas:

- Carregue a base de dados `media_precos_carros_brasil.csv`;
- Verifique se há valores faltantes nos dados. Caso haja, escolha uma tratativa para resolver o problema de valores faltantes;
- Verifique se há dados duplicados nos dados;

- d) Crie duas categorias, para separar colunas numéricas e categóricas. Imprima o resumo de informações das variáveis numéricas e categóricas (estatística descritiva dos dados);
- e) Imprima a contagem de valores por modelo (`model`) e marca do carro (`brand`);
- f) Dê uma breve explicação (máximo de quatro linhas) sobre os principais resultados encontrados na Análise Exploratória dos dados.

2 Visualização dos Dados

A partir da base de dados `precos_carros_brasil.csv`, execute as seguintes tarefas:

- a) Gere um gráfico da distribuição da quantidade de carros por marca;
- b) Gere um gráfico da distribuição da quantidade de carros por tipo de engrenagem do carro;
- c) Gere um gráfico da evolução da média de preço dos carros ao longo dos meses de 2022 (variável de tempo no eixo X);
- d) Gere um gráfico da distribuição da média de preço dos carros por marca e tipo de engrenagem;
- e) Dê uma breve explicação (máximo de quatro linhas) sobre os resultados gerados no item d;
- f) Gere um gráfico da distribuição da média de preço dos carros por marca e tipo de combustível;
- g) Dê uma breve explicação (máximo de quatro linhas) sobre os resultados gerados no item f.

3 Aplicação de Modelos de *Machine Learning* para Prever o Preço Médio dos Carros

A partir da base de dados `precos_carros_brasil.csv`, execute as seguintes tarefas:

- a) Escolha as variáveis **numéricas** (modelos de Regressão) para serem as variáveis independentes do modelo. A variável *target* é `avg_price`.
Observação: caso julgue necessário, faça a transformação de variáveis categóricas em variáveis numéricas para inputar no modelo. Indique foram transformadas e **como** foram transformadas;
- b) Crie partições contendo 75% dos dados para treino e 25% para teste;
- c) Treine modelos `RandomForestRegressor` e `XGBRegressor` para predição dos preços dos carros. Observação: caso julgue necessário, mude os parâmetros dos modelos e rode novos modelos. Indique quais parâmetros foram inputados e indique o treinamento de cada modelo;
- d) Grave os valores preditos em variáveis criadas;

- e) Realize a análise de importância das variáveis para estimar a variável *target*, para cada modelo treinado;
- f) Dê uma breve explicação (máximo de quatro linhas) sobre os resultados encontrados na análise de importância de variáveis;
- g) Escolha o melhor modelo com base nas métricas de avaliação MSE, MAE e R^2 ;
- h) Dê uma breve explicação (máximo de quatro linhas) sobre qual modelo gerou o melhor resultado e a métrica de avaliação utilizada.

B - RESOLUÇÃO

1 Análise Exploratória dos Dados

a. Carregue a base de dados `media_precos_carros_brasil.csv`

```

1 # Importando bibliotecas necessárias
2 import pandas as pd
3 import matplotlib.pyplot as plt
4 import seaborn as sns
5 import warnings
6 from sklearn.preprocessing import LabelEncoder
7 from sklearn.model_selection import train_test_split
8 from sklearn.ensemble import RandomForestRegressor
9 from xgboost import XGBRegressor
10 from sklearn.metrics import mean_squared_error, mean_absolute_error,
    r2_score
11
12 warnings.filterwarnings('ignore')
13
14 # Carregando dados do arquivo precos_carros_brasil.csv
15 dados = pd.read_csv('precos_carros_brasil.csv')
16
17 # Listando o nome das colunas
18 dados.columns
19
20 # Imprimindo somente as cinco primeiras linhas
21 dados.head()
22
23 # Imprime o tipo de dado de cada coluna
24 dados.dtypes
25
26 # Número de linhas e colunas
27 dados.shape

```

b. Verifique se há valores faltantes nos dados. Caso haja, escolha uma tratativa para resolver o problema de valores faltantes

```

1 # Verificando se há valores faltantes nos dados
2 dados.isna().any()
3 dados.isna().sum()
4
5 # Verificando se há linhas inteiramente vazias e quantas existem

```

```

6 dados.isnull().all(axis=1).sum()
7
8 # Removendo todas as linhas que têm todos os campos vazios
9 dados.dropna(axis=0, how='all', inplace=True)
10
11 # Verificando novamente se há valores faltantes
12 dados.isna().sum()

```

c. Verifique se há dados duplicados nos dados

```

1 # Verificando se há dados duplicados
2 dados.duplicated().sum()
3
4 # Removendo dados duplicados
5 dados.drop_duplicates(inplace=True)

```

d. Crie duas categorias, para separar colunas numéricas e categóricas. Imprima o resumo de informações das variáveis numéricas e categóricas (estatística descritiva dos dados)

```

1 # Convertendo colunas numéricas para tipos adequados
2 dados['engine_size'] = dados['engine_size'].str.replace(',', '.').
   astype(float)
3 dados['year_of_reference'] = pd.to_numeric(dados['year_of_reference']
   ], errors='coerce').astype('Int64')
4 dados['year_model'] = pd.to_numeric(dados['year_model'], errors='
   coerce').astype('Int64')
5 dados['avg_price_brl'] = pd.to_numeric(dados['avg_price_brl'], errors
   ='coerce')
6
7 # Criando duas categorias: numéricas e categóricas
8 numericas_cols = [col for col in dados.columns if dados[col].dtype !=
   'object']
9
10 categoricas_cols = [col for col in dados.columns if dados[col].dtype
   == 'object']
11
12 # Estatística descritiva das variáveis numéricas
13 dados[numericas_cols].describe()
14
15 # Estatística descritiva das variáveis categóricas
16 dados[categoricas_cols].describe()

```

e. Imprima a contagem de valores por modelo (model) e marca do carro (brand)

```

1 # Contagem de valores por modelo e marca
2 dados.groupby(['model', 'brand']).size().reset_index(name='contagem')

```

f) Dê uma breve explicação (máximo de quatro linhas) sobre os principais resultados encontrados na Análise Exploratória dos dados.

No geral, o preço médio de um carro no Brasil custa aproximadamente R\$ 52.756,91. Um dos modelos mais vendidos é Palio Week. Adv/Adv TRYON 1.8 mpi Flex, com câmbio manual e motor 1.6. O menor preço médio de carros foi de R\$6.647,00, e o maior preço médio é R\$ 979.358,00.

2 Visualização dos Dados

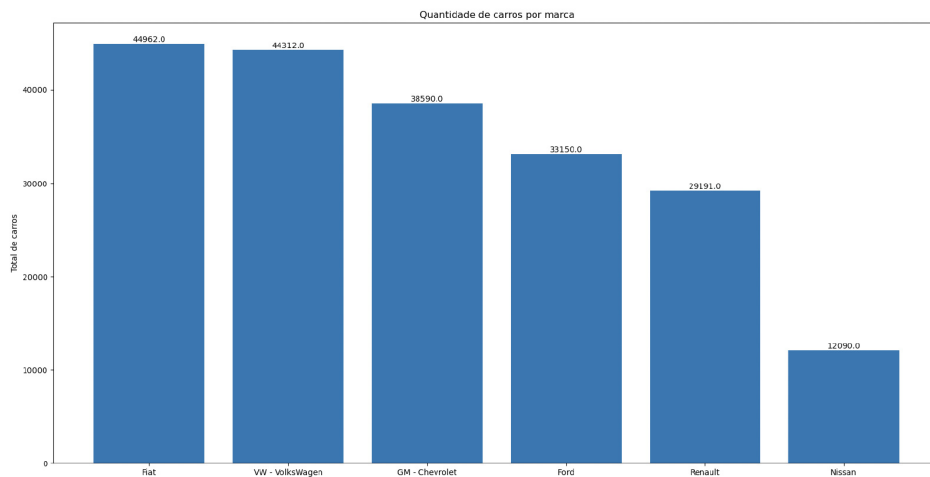
a. Gere um gráfico da distribuição da quantidade de carros por marca

```

1 # Gráfico: distribuição da quantidade de carros por marca
2 plt.figure(figsize=(20,10))
3 grafic_brands_qties = plt.bar(dados['brand'].value_counts().index,
4                               dados['brand'].value_counts().values)
5 plt.title('Quantidade de carros por marca')
6 plt.ylabel('Total de carros')
7 plt.bar_label(grafic_brands_qties, size=10, fmt="%.01f", label_type='
  edge')

```

Figura 2.1 – RESOLUÇÃO EXERCÍCIO 2: DISTRIBUIÇÃO DA QUANTIDADE DE CARROS POR MARCA



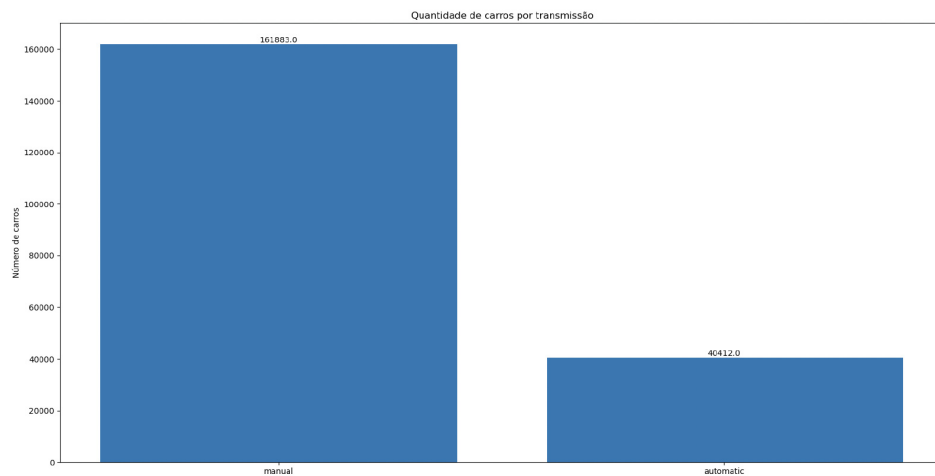
b. Gere um gráfico da distribuição da quantidade de carros por tipo de engrenagem do carro

```

1 # Gráfico: distribuição da quantidade de carros por tipo de
  engrenagem
2 plt.figure(figsize=(20,10))
3 grafic_gear_qties = plt.bar(dados['gear'].value_counts().index,
4                              dados['gear'].value_counts().values)
5 plt.title('Quantidade de carros por transmissão')
6 plt.ylabel('Número de carros')
7 plt.bar_label(grafic_gear_qties, size=10, fmt="%.01f", label_type='
  edge')

```

Figura 2.2 – RESOLUÇÃO EXERCÍCIO 2: DISTRIBUIÇÃO DA QUANTIDADE DE CARROS POR TIPO DE ENGRENAGEM

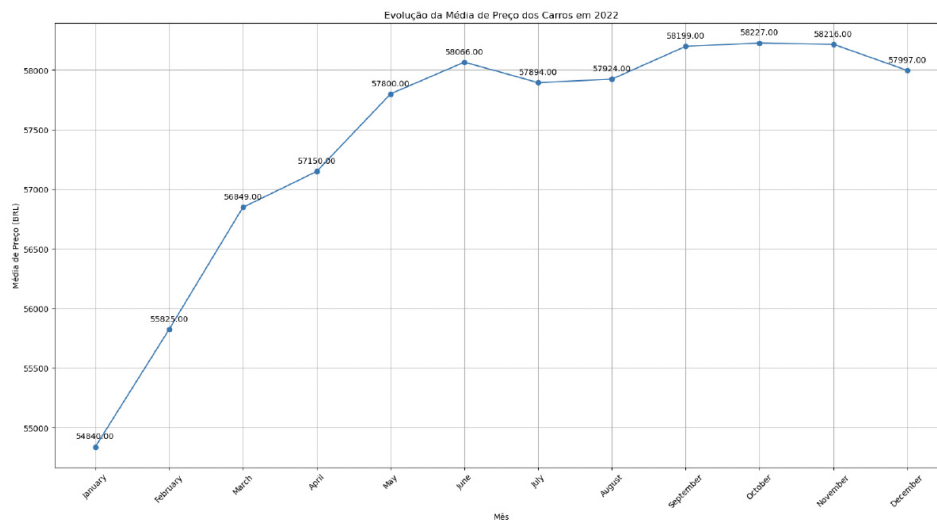


c. Gere um gráfico da evolução da média de preço dos carros ao longo dos meses de 2022 (variável de tempo no eixo X)

```

1 # Evolução da média de preço dos carros ao longo de 2022
2 dados['year_of_reference'] = dados['year_of_reference'].round().
  astype(int)
3 cars_in_2022 = dados[dados['year_of_reference'] == 2022]
4 month_order = ['January', 'February', 'March', 'April', 'May', 'June', '
  July',
5               'August', 'September', 'October', 'November', 'December']
6 cars_in_2022['month_of_reference'] = pd.Categorical(cars_in_2022['
  month_of_reference'],
7                                                     categories=month_order, ordered
  =True)
8 cars_avg_price_in_2022 = cars_in_2022.groupby(['month_of_reference'])
  ['avg_price_brl'].mean().round(0).reset_index(name='average_price'
  )
9
10 plt.figure(figsize=(20,10))
11 plt.plot(cars_avg_price_in_2022['month_of_reference'],
12          cars_avg_price_in_2022['average_price'], marker='o', linestyle
  ='-')
13 for x, y in zip(cars_avg_price_in_2022['month_of_reference'],
14                cars_avg_price_in_2022['average_price']):
15     plt.annotate(f"{y:.2f}", (x,y), textcoords="offset points", xytext
16                 =(0,10), ha='center')
17 plt.title('Evolução da Média de Preço dos Carros em 2022')
18 plt.xlabel('Mês')
19 plt.ylabel('Média de Preço (BRL)')
20 plt.xticks(cars_avg_price_in_2022.index, month_order, rotation=45)
21 plt.grid(True)
  
```

Figura 2.3 – RESOLUÇÃO EXERCÍCIO 2: EVOLUÇÃO DA MÉDIA DE PREÇO DOS CARROS AO LONGO DE 2022



d. Gere um gráfico da distribuição da média de preço dos carros por marca e tipo de engrenagem

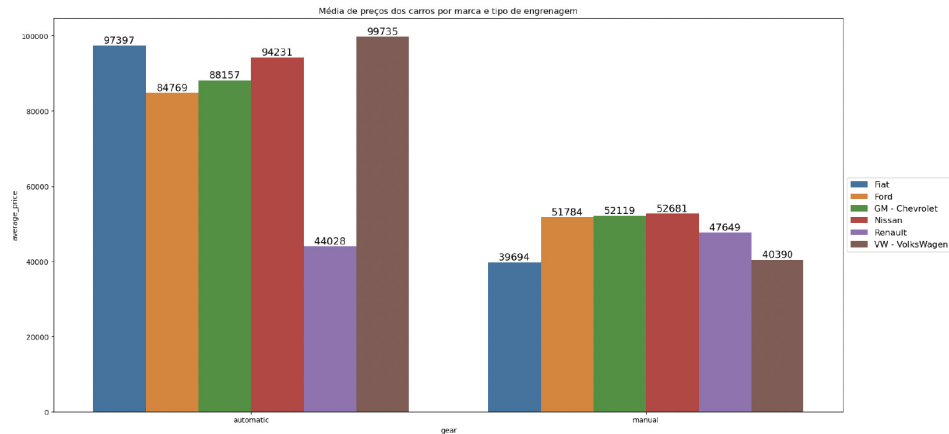
```

1 # Agrupar média de preços médios dos carros por marca e tipo de
  engrenagem
2 car_gear_brand_avg_price = dados.groupby(['brand', 'gear'])['
  avg_price_brl'].mean().round(0)
3 car_gear_brand_avg_price.head(12)
4
5 # Utilizando a função reset_index para criar uma ordem e facilitar a
  criação do gráfico
6 car_gear_brand_avg_price = car_gear_brand_avg_price.reset_index(name=
  'average_price')
7 car_gear_brand_avg_price.head(12)
8
9 # Gerar o gráfico de distribuição média dos preços dos carros por
  marca e tipo de engrenagem
10 plt.figure(figsize=(20,10))
11 barplot = sns.barplot(
12 x='gear',
13 y='average_price',
14 hue='brand',
15 data=car_gear_brand_avg_price,
16 hue_order=car_gear_brand_avg_price['brand'].unique(),
17 )
18
19 barplot.legend(loc='center left', bbox_to_anchor=(1, 0.5), ncol=1,
  fontsize=12)
20 barplot.bar_label(barplot.containers[0], fontsize=14);
21 barplot.bar_label(barplot.containers[1], fontsize=14);
22 barplot.bar_label(barplot.containers[2], fontsize=14);
23 barplot.bar_label(barplot.containers[3], fontsize=14);
24 barplot.bar_label(barplot.containers[4], fontsize=14);
25 barplot.bar_label(barplot.containers[5], fontsize=14);

```

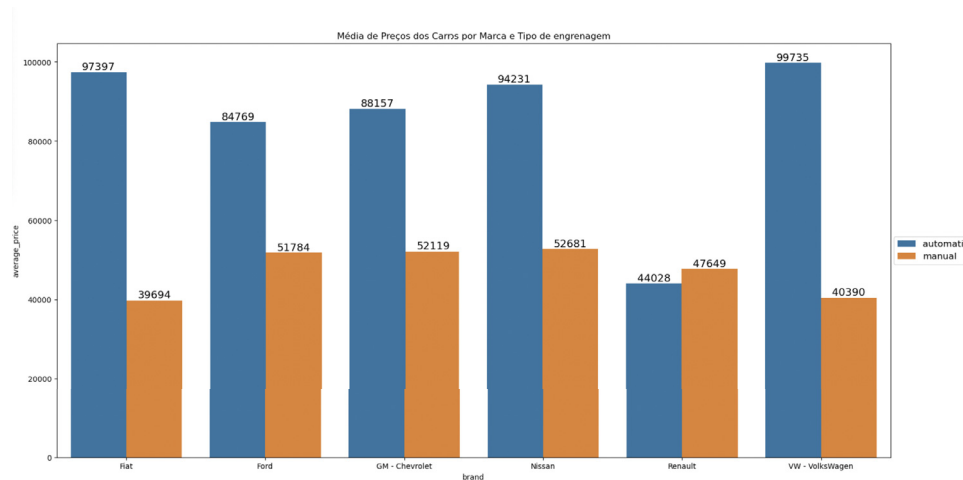
```
26 plt.title('Média de preços dos carros por marca e tipo de engrenagem')
    )
```

Figura 2.4 – RESOLUÇÃO EXERCÍCIO 2: MÉDIA DE PREÇOS DOS CARROS POR MARCA E TIPO DE ENGRENAGEM AGRUPADOS POR TIPO



```
1 # Gerar o segundo tipo de gráfico de distribuição média dos preços
  dos carros por marca e tipo de engrenagem
2 plt.figure(figsize=(20,10))
3 barplot = sns.barplot(
4 x='brand',
5 y='average_price',
6 hue='gear',
7 data=car_gear_brand_avg_price,
8 hue_order=car_gear_brand_avg_price['gear'].unique(),)
9 barplot.legend(loc='center left', bbox_to_anchor=(1, 0.5), ncol=1,
  fontsize=12)
10 barplot.bar_label(barplot.containers[0], fontsize=14);
11 barplot.bar_label(barplot.containers[1], fontsize=14);
12 plt.title('Média de Preços dos Carros por Marca e Tipo de engrenagem')
    )
```

Figura 2.5 – RESOLUÇÃO EXERCÍCIO 2: MÉDIA DE PREÇOS DOS CARROS POR MARCA E TIPO DE ENGRENAGEM AGRUPADOS POR MARCA



e. Dê uma breve explicação (máximo de quatro linhas) sobre os resultados gerados no item d:

O preço médio do carro com câmbio automático é relativamente mais alto quando comparado aos de câmbio manual para todas as marcas, com exceção da Renault. Volkswagen e Fiat lideram as marcas que possuem maior preço médio para carros de câmbio automático, ao passo que lideram também as marcas com menor preço médio dos carros com câmbio manual.

f. Gere um gráfico da distribuição da média de preço dos carros por marca e tipo de combustível

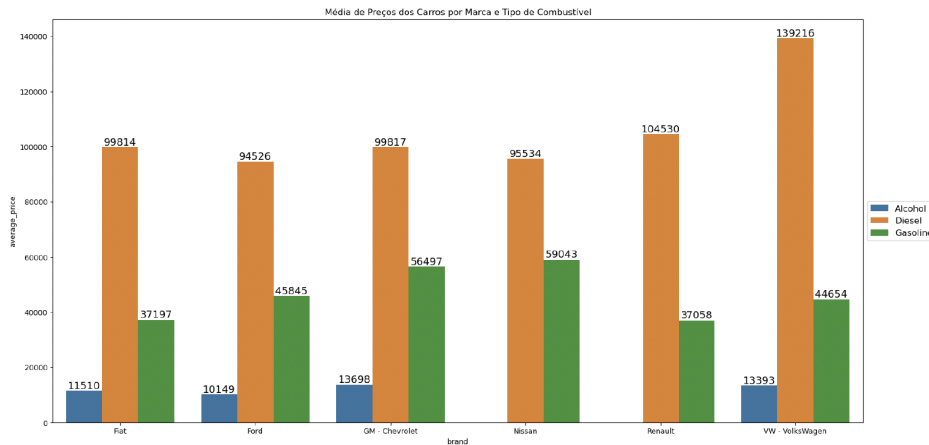
```

1 # Agrupar média de preços médios dos carros por marcar e tipo de
  combustivel
2 cars_avg_price_brand_fuel = dados.groupby(['brand', 'fuel'])['
  avg_price_brl'].mean().round(0)
3 cars_avg_price_brand_fuel.head(16)
4
5 # Utilizando a função reset_index para criar uma ordem e facilitar a
  criação do gráfico
6 cars_avg_price_brand_fuel = cars_avg_price_brand_fuel.reset_index(
  name='average_price')
7 cars_avg_price_brand_fuel.head(16)
8
9 # Gerar gráfico de média de preços médios de carros por marcar e tipo
  de combustivel
10 plt.figure(figsize=(20, 10))
11 barplot = sns.barplot(
12 x='brand',
13 y='average_price',
14 hue='fuel',
15 data=cars_avg_price_brand_fuel,
16 hue_order=cars_avg_price_brand_fuel['fuel'].unique()
17 )
18 barplot.legend(loc='center left', bbox_to_anchor=(1, 0.5), ncol=1,
  fontsize=12)
19 barplot.bar_label(barplot.containers[0], fontsize=14)
20 barplot.bar_label(barplot.containers[1], fontsize=14)
21 barplot.bar_label(barplot.containers[2], fontsize=14)

```

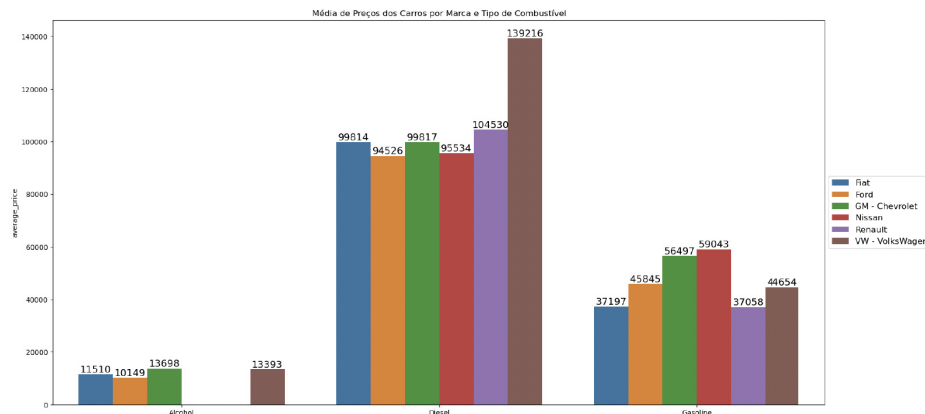
```
plt.title('Média de Preços dos Carros por Marca e Tipo de Combustível')
```

Figura 2.6 – RESOLUÇÃO EXERCÍCIO 2: MÉDIA DE PREÇOS DOS CARROS POR MARCA E TIPO DE COMBUSTÍVEL AGRUPADOS POR MARCA



```
1 # Gerar gráfico de média de preços médios de carros por marcar e tipo
  de combustível - Segundo tipo
2 plt.figure(figsize=(20, 10))
3 barplot = sns.barplot(
4 x='fuel',
5 y='average_price',
6 hue='brand',
7 data=cars_avg_price_brand_fuel,
8 hue_order=cars_avg_price_brand_fuel['brand'].unique())
9
10 barplot.legend(loc='center left', bbox_to_anchor=(1, 0.5), ncol=1,
11               fontsize=12)
12 barplot.bar_label(barplot.containers[0], fontsize=14)
13 barplot.bar_label(barplot.containers[1], fontsize=14)
14 barplot.bar_label(barplot.containers[2], fontsize=14)
15 barplot.bar_label(barplot.containers[3], fontsize=14)
16 barplot.bar_label(barplot.containers[4], fontsize=14)
17 barplot.bar_label(barplot.containers[5], fontsize=14)
18 plt.title('Média de Preços dos Carros por Marca e Tipo de Combustível')
```

Figura 2.7 – RESOLUÇÃO EXERCÍCIO 2: MÉDIA DE PREÇOS DOS CARROS POR MARCA E TIPO DE COMBUSTÍVEL AGRUPADOS POR TIPO



g. Dê uma breve explicação (máximo de quatro linhas) sobre os resultados gerados no item f:

1. O preço médio do carro à diesel é muito mais elevado para todas as marcas, sendo a Volkswagen a marca com maior preço médio para esta categoria.
2. O preço médio do carro à álcool lidera com os menores valores.
3. Os carros à gasolina com menor preço médio são das marcas Fiat e Renault.

3 Aplicação de Modelos de Machine Learning para Prever o Preço Médio dos Carros

a) Escolha as variáveis numéricas (modelos de Regressão) para serem as variáveis independentes do modelo. A variável target é avg_price.

```

1 # Gráfico de distribuição boxplot dos preços médios dos carros
2 sns.boxplot(dados['avg_price_brl']).set_title('Distribuição dos preços médios')
3
4 # Transformar mês de referência em uma variável numérica
5 dados['month_of_reference'] = LabelEncoder().fit_transform(dados['month_of_reference'])
6
7 # Transformar marca de carro em uma variável numérica
8 dados['brand'] = LabelEncoder().fit_transform(dados['brand'])
9
10 # Transformar modelo de carro em uma variável numérica
11 dados['model'] = LabelEncoder().fit_transform(dados['model'])
12
13 # Transformar tipo de combustível em uma variável numérica
14 dados['fuel'] = LabelEncoder().fit_transform(dados['fuel'])
15
16 # Transformar tipo de transmissão em uma variável numérica
17 dados['gear'] = LabelEncoder().fit_transform(dados['gear'])
18
19 # Remover as variáveis de entrada que não são consideradas importantes para a predição
20 dados_interesting = dados.drop([

```

```

21 'fipecode',
22 'authentication',
23 ], axis=1)
24
25 # Mapa de correlação das variáveis numéricas com variável Target
26 sns.heatmap(dados_interesting.corr("spearman"), annot = True)
27 plt.title("Mapa de Correlação das Variáveis Numéricas\n", fontsize =
28           18)
29 plt.show()
30
31 # Variável X contém apenas variáveis numéricas de interesse para a an
32   álise excluindo a variável target
33
34 X = dados_interesting.drop(['avg_price_brl'], axis=1)
35 X.head()
36
37 # Variável Y contém apenas a variável target - avg_price_brl
38 Y = dados_interesting['avg_price_brl']
39 Y.head()

```

Figura 2.8 – RESOLUÇÃO EXERCÍCIO 3: BOXPLOT DA DISTRIBUIÇÃO DOS PREÇOS MÉDIOS

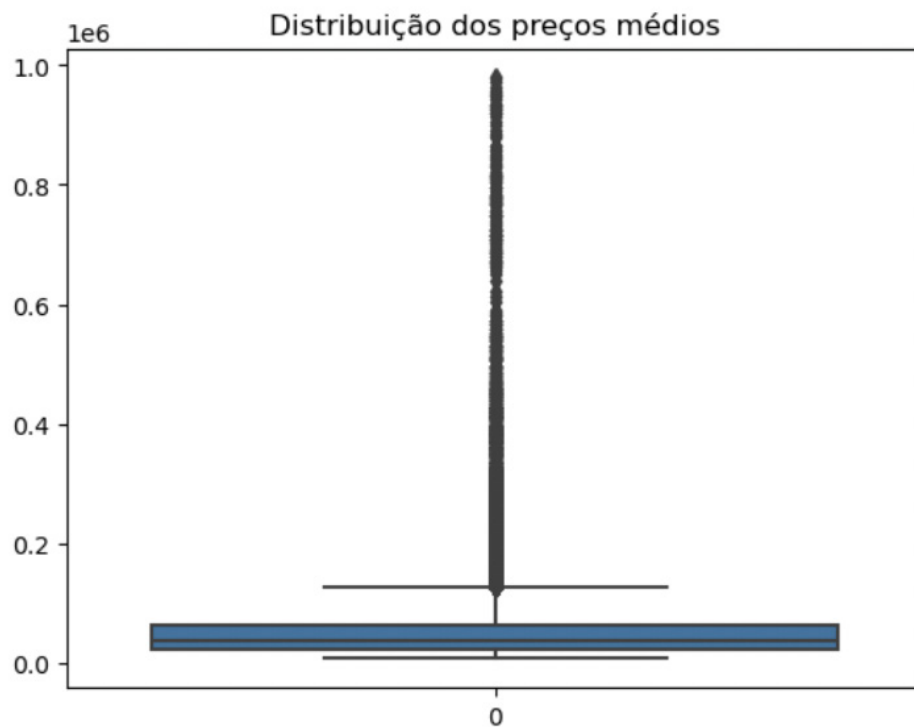
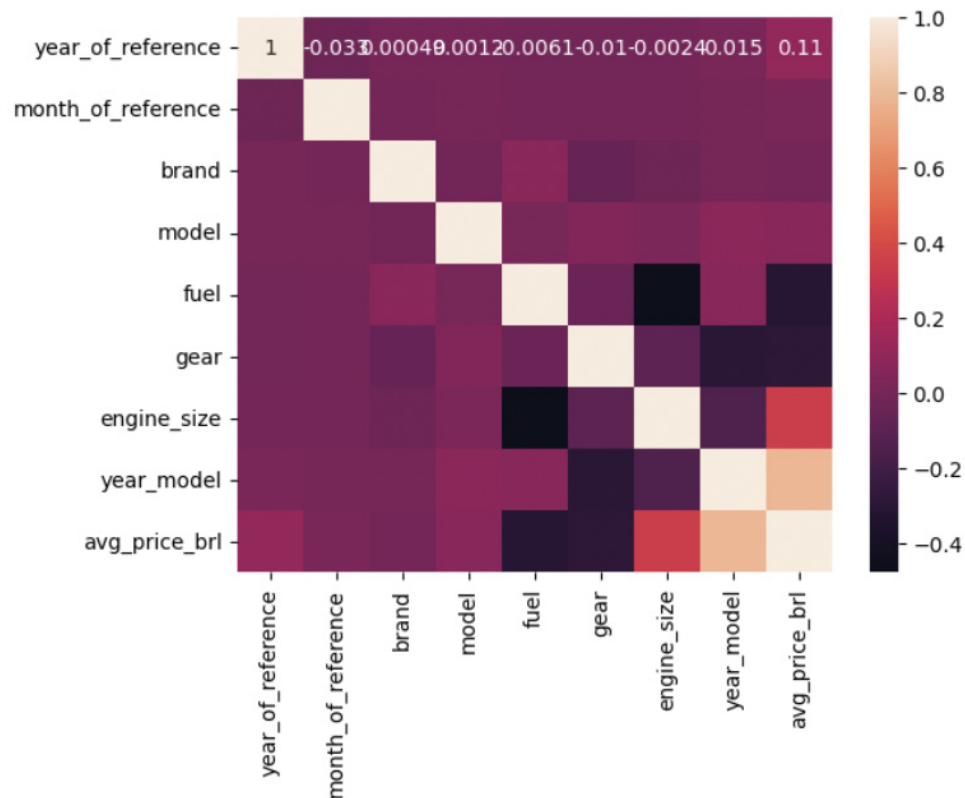


Figura 2.9 – RESOLUÇÃO EXERCÍCIO 3: MAPA DE CORRELAÇÃO ENTRE AS VARIÁVEIS NUMÉRICAS



b. Crie partições contendo 75% dos dados para treino e 25% para teste

```
1 # Divisão: 25% dos dados são de teste e 75% de treinamento
2 X_train, x_test, Y_train, y_test = train_test_split(X, Y, test_size =
    0.25, random_state = 42)
```

c. Treine modelos RandomForest (biblioteca RandomForestRegressor) e XGBoost (biblioteca XGBRegressor) para predição dos preços dos carros

```
1 #Algoritmo Random Forest, sem especificar nenhum parâmetro (número de
    árvores, número de ramificações, etc)
2 model_rf = RandomForestRegressor()
3
4 # Random Forest com os hiperparâmetros:
5 model_rf_paramenters = RandomForestRegressor(
6 max_depth=29,
7 min_samples_leaf=32,
8 min_samples_split=28,
9 n_estimators=208,
10 random_state=43
11 )
12 # Ajuste do modelo, de acordo com as variáveis de treinamento
13 model_rf_paramenters.fit(X_train, Y_train)
14
15 # Ajuste do modelo, de acordo com as variáveis de treinamento
16 model_rf.fit(X_train, Y_train)
17
18
```

```

19 # Algoritmo XGBoost sem parametros
20 model_xg = XGBRegressor()
21
22 # Ajuste do modelo, de acordo com as variáveis de treinamento
23 model_xg.fit(X_train, Y_train)
24
25 #Algoritmo XGBoost com parametros
26 model_parameters_xg = XGBRegressor(
27     learning_rate=1,
28     n_estimators=100,
29     max_depth=25,
30     min_child_weight=1,
31     gamma=0,
32     subsample=0.8,
33     colsample_bytree=0.8,
34     reg_alpha=0,
35     reg_lambda=1,
36     random_state=42,
37     n_jobs=-1
38 )
39
40 # Ajuste do modelo, de acordo com as variáveis de treinamento
41 model_parameters_xg.fit(X_train, Y_train)

```

d. Grave os valores preditos em variáveis criadas

```

1 # Predição dos valores de preço médio dos carros com base nos dados
  de teste Random Forest sem parametros
2 predicted_values_rf = model_rf.predict(x_test)
3
4 # Predição dos valores de preço médio dos carros com base nos dados
  de teste Random Forest com parametros
5 predicted_values_parameters_rf = model_rf_parameters.predict(x_test)
6
7 # Predição dos valores de preços médios dos carros com base nos dados
  de teste XGBoost sem parametros
8 predicted_values_xg = model_xg.predict(x_test)
9
10 # Predição dos valores de preços médios dos carros com base nos dados
   de teste XGBoost com parametros
11 predicted_values_parameters_xg = model_parameters_xg.predict(x_test)

```

e. Realize a análise de importância das variáveis para estimar a variável target, para cada modelo treinado

```

1 # Análise da importancia das variáveis para estimar a variável target
  - Random Forest sem parametros
2 model_rf.feature_importances_
3 feature_importances_rf = pd.DataFrame(
4 model_rf.feature_importances_,
5 index = X_train.columns,
6 columns = ['importances']
7 ).sort_values('importances', ascending=True)

```

month_of_reference 0.005168
 year_of_reference 0.012553
 brand 0.016695
 fuel 0.032713
 gear 0.035363
 model 0.058640
 year_model 0.388397
 engine_size 0.450472

```

1 # Análise da importancia das variáveis para estimar a variável target
  - Random Forest com parametros
2 model_rf_paramenters.feature_importances_
3 feature_importances_parameters_rf = pd.DataFrame(
4 model_rf_paramenters.feature_importances_,
5 index = X_train.columns,
6 columns = ['importances']
7 ).sort_values('importances', ascending=True)
8 feature_importances_parameters_rf
  
```

month_of_reference 0.000690
 year_of_reference 0.010383
 brand 0.016098
 gear 0.021959
 fuel 0.034132
 model 0.037028
 year_model 0.410473
 engine_size 0.469236

```

1 # Análise da importancia das variáveis para estimar a variável target
  - XGBoost com parametros
2 model_parameters_xg.feature_importances_
3 feature_importances_parameters_xg = pd.DataFrame(
4 model_rf_paramenters.feature_importances_,
5 index = X_train.columns,
6 columns = ['importances']
7 ).sort_values('importances', ascending=True)
8 feature_importances_parameters_xg
  
```

month_of_reference 0.000690
 year_of_reference 0.010383
 brand 0.016098
 gear 0.021959
 fuel 0.034132
 model 0.037028
 year_model 0.410473
 engine_size 0.469236

```

1 # Analisando a importância das variáveis para estimar a variável
  target - XGBoost sem parametros
2 model_xg.feature_importances_
3 feature_importances_xg = pd.DataFrame(
4 model_xg.feature_importances_,
  
```

```

5 index=X_train.columns,
6 columns=['importances']
7 ).sort_values('importances', ascending=True)
8 feature_importances_xg

```

```

month_of_reference 0.004271
year_of_reference 0.017475
model 0.023331
brand 0.056276
gear 0.122657
fuel 0.154437
year_model 0.190992
engine_size 0.430561

```

f. Dê uma breve explicação (máximo de quatro linhas) sobre os resultados encontrados na análise de importância de variáveis:

As variáveis que tiveram maior relevância na medida da performance dos modelos treinados foram o tamanho do motor (engine_size) e ano do modelo (year_model) respectivamente. O mês de referência foi o que obteve menor relevância em todos o modelos treinados.

g. Escolha o melhor modelo com base nas métricas de avaliação MSE, MAE e R2:

```

1 # MSE - calcula o erro quadrático médio das predições do modelo
2 # Random Forest sem parametro
3 mse_rf = mean_squared_error(y_test, predicted_values_rf) #12181571.08
4
5 # O MAE calcula a média da diferença absoluta entre o valor predito e
  o valor real
6 # Random Forest sem parametro
7 mea_rf = mean_absolute_error(y_test, predicted_values_rf) # 1765.88
8
9 # O R2 é uma métrica que varia entre 0 e 1 e é uma razão que indica o
  quão bom o nosso modelo
10 # Random Forest sem parametros
11 r2_score(y_test, predicted_values_rf) # 0.995

```

```

1 # O MSE - Random Forest com parametros
2 mse_parameters_rf = mean_squared_error(y_test,
  predicted_values_parameters_rf) # 151743671.60
3
4 # O MAE - Random Forest com parametros
5 mae_parameters_rf = mean_absolute_error(y_test,
  predicted_values_parameters_rf) # 4577.49
6
7 # R2 score - Random Forest com parametros # 0.943
8 r2_score(y_test, predicted_values_parameters_rf)

```

```

1 # MSE - XGBoost sem parametros
2 mse_xg = mean_squared_error(y_test, predicted_values_xg) #
  30952404.229
3
4 # MAE - XGBoost sem parametros
5 mae_xg = mean_absolute_error(y_test, predicted_values_xg) # 3245.31

```

```

6
7 # R2 score - XGBoost sem parametros
8 r2_score(y_test, predicted_values_xg) # 0.988

1 # MSE - XGBoost com parametros
2 mse_parameters_xg = mean_squared_error(y_test,
    predicted_values_parameters_xg) # 58122007.14
3
4 # MAE - XGBoost com parametros
5 mae_parameters_xg = mean_absolute_error(y_test,
    predicted_values_parameters_xg) # 3988.81
6
7 # R2 score - XGBoost com parametros
8 r2_score(y_test, predicted_values_parameters_xg) # 0.978

```

O melhor modelo com base nessas medidas de acurácia foi o modelo RandomForest sem parâmetros.

h. Dê uma breve explicação (máximo de quatro linhas) sobre qual modelo gerou o melhor resultado e a métrica de avaliação utilizada:

O modelo que apresentou os melhores resultado foi o RandomForest sem parametros para medir a sua acurácia foi utilizado as medidas de acurácia MSE (*mean squared error*) com o valor 12165425.33231638, MAE (*mean absolute error*) com o valor 12165425.33 e R2_score com 0.995 de precisão.

APÊNDICE 3 – LINGUAGEM R

A - ENUNCIADO

1 Pesquisa com Dados de Satélite (Satellite)

O banco de dados consiste nos valores multiespectrais de pixels em vizinhanças 3x3 em uma imagem de satélite, e na classificação associada ao pixel central em cada vizinhança. O objetivo é prever esta classificação, dados os valores multiespectrais.

Um quadro de imagens do Satélite Landsat com MSS (Multispectral Scanner System) consiste em quatro imagens digitais da mesma cena em diferentes bandas espectrais. Duas delas estão na região visível (correspondendo aproximadamente às regiões verde e vermelha do espectro visível) e duas no infravermelho (próximo). Cada pixel é uma palavra binária de 8 bits, com 0 correspondendo a preto e 255 a branco. A resolução espacial de um pixel é de cerca de 80m x 80m. Cada imagem contém 2340 x 3380 desses pixels. O banco de dados é uma subárea (minúscula) de uma cena, consistindo de 82 x 100 pixels. Cada linha de dados corresponde a uma vizinhança quadrada de pixels 3x3 completamente contida dentro da subárea 82x100. Cada linha contém os valores de pixel nas quatro bandas espectrais (convertidas em ASCII) de cada um dos 9 pixels na vizinhança de 3x3 e um número indicando o rótulo de classificação do pixel central.

As classes são: solo vermelho, colheita de algodão, solo cinza, solo cinza úmido, restolho de vegetação, solo cinza muito úmido.

Os dados estão em ordem aleatória e certas linhas de dados foram removidas, portanto você não pode reconstruir a imagem original desse conjunto de dados. Em cada linha de dados, os quatro valores espectrais para o pixel superior esquerdo são dados primeiro, seguidos pelos quatro valores espectrais para o pixel superior central e, em seguida, para o pixel superior direito, e assim por diante, com os pixels lidos em sequência, da esquerda para a direita e de cima para baixo. Assim, os quatro valores espectrais para o pixel central são dados pelos atributos 17, 18, 19 e 20. Se você quiser, pode usar apenas esses quatro atributos, ignorando os outros. Isso evita o problema que surge quando uma vizinhança 3x3 atravessa um limite.

O banco de dados se encontra no pacote mlbench e é completo (não possui dados faltantes).

Tarefas:

1. Carregue a base de dados Satellite.
2. Crie partições contendo 80% para treino e 20% para teste.
3. Treine modelos RandomForest, SVM e RNA para predição destes dados.
4. Escolha o melhor modelo com base em suas matrizes de confusão.
5. Indique qual modelo dá o melhor o resultado e a métrica utilizada.

2 Estimativa de Volumes de Árvores

Modelos de aprendizado de máquina são bastante usados na área da engenharia florestal (mensuração florestal) para, por exemplo, estimar o volume de madeira de árvores sem ser necessário abatê-las.

O processo é feito pela coleta de dados (dados observados) através do abate de algumas árvores, onde sua altura, diâmetro na altura do peito (dap), etc, são medidos de forma exata.

Com estes dados, treina-se um modelo de AM que pode estimar o volume de outras árvores da população.

Os modelos, chamados alométricos, são usados na área há muitos anos e são baseados em regressão (linear ou não) para encontrar uma equação que descreve os dados. Por exemplo, o modelo de Spurr é dado por:

$$Volume = b0 + b1 * dap2 * Ht \quad (3.1)$$

Onde dap é o diâmetro na altura do peito (1,3 metros), Ht é a altura total. Tem-se vários modelos alométricos, cada um com uma determinada característica, parâmetros, etc. Um modelo de regressão envolve aplicar os dados observados e encontrar b0 e b1 no modelo apresentado, gerando assim uma equação que pode ser usada para prever o volume de outras árvores.

Dado o arquivo Volumes.csv, que contém os dados de observação, escolha um modelo de aprendizado de máquina com a melhor estimativa, a partir da estatística de correlação.

Tarefas

1. Carregar o arquivo Volumes.csv (<http://www.razer.net.br/datasets/Volumes.csv>)
2. Eliminar a coluna NR, que só apresenta um número sequencial
3. Criar partição de dados: treinamento 80%, teste 20%
4. Usando o pacote "caret", treinar os modelos: Random Forest (rf), SVM (svmRadial), Redes Neurais (neuralnet) e o modelo alométrico de SPURR. O modelo alométrico é dado por:

$$Volume = b0 + b1 * dap2 * Ht \quad (3.2)$$

```
1 alom <- nls(VOL ~ b0 + b1*DAP*DAP*HT, dados, start=list(b0
  =0.5, b1=0.5))
```

5. Efetue as predições nos dados de teste
6. Crie suas próprias funções (UDF) e calcule as seguintes métricas entre a predição e os dados observados
 - Coeficiente de determinação : R^2

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (3.3)$$

onde y_i é o valor observado, $\hat{y}_i$ é o valor predito e $\bar{y}$ é a média dos valores y_i observados. Quanto mais perto de 1 melhor é o modelo;

- Erro padrão da estimativa: S_{yx}

$$S_{yx} = \sqrt{\frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{n - 2}} \quad (3.4)$$

esta métrica indica erro, portanto quanto mais perto de 0 melhor é o modelo;

- $Sy\%$:

$$S_{yx}\% = \frac{S_{yx}}{\bar{y}} \times 100 \quad (3.5)$$

esta métrica indica porcentagem de erro, portanto quanto mais perto de 0 melhor é o modelo;

7. Escolha o melhor modelo.

B - RESOLUÇÃO

1 Pesquisa com Dados de Satélite (Satellite)

```

1 ##### 1.1 Carregue a base de dados Satellite
2
3 #instalar o pacotes necessários
4 #install.packages("mlbench", repos = "http://cran.us.r-project.org")
5 #install.packages("e1017", repos = "http://cran.us.r-project.org")
6 #install.packages("randomForest", repos = "http://cran.us.r-project.
  org")
7 #install.packages("kernlab", repos = "http://cran.us.r-project.org")
8 #install.packages("caret", repos = "http://cran.us.r-project.org")
9
10 # usar os pacotes necessários
11 library(mlbench)
12 library(caret)
13
14 # carregar os dados Satellite
15 data(Satellite)
16
17 # exibir estrutura dos dados Satellite
18 str(Satellite)
19
20 # apresentar alguma medidas estatísticas do dados Satellite
21 summary(Satellite)
22
23 # exibir alguns dados do Satellite
24 head(Satellite, n = 6)
25
26 ##### 1.2 Crie partições contendo 80% para treino e 20% para teste
27
28 # Para reproductibilidade
29 set.seed(7)
30
31 # particionar em 80% para treino e 20% para teste
32 indices <- createDataPartition(Satellite[["classes"]], p=0.8, list=F)
33 treino <- Satellite[indices, ]
34 teste <- Satellite[-indices, ]
35
36 ##### 1.3 Treine modelos RandomForest, SVM e RNA para predição destes
  dados.
37
38 # treinar modelos RandomForest, SVM e RNA

```

```

39 rf <- train(classes ~ ., data=treino, method="rf")
40 svm <- train(classes ~ ., data=treino, method="svmRadial")
41 rna <- train(classes ~ ., data=treino, method="nnet", trace=F)
42
43 ##### 1.4 Escolha o melhor modelo com base em suas matrizes de confusã
    o.
44
45 # previsões
46 predict.rf <- predict(rf, teste)
47 predict.svm <- predict(svm, teste)
48 predict.rna <- predict(rna, teste)
49
50 # matrizes de confusões de cada uma das previsões
51
52 # matriz de confusão para o modelo RF
53 conf_matrix.rf <- confusionMatrix(predict.rf, teste[["classes"]])
54
55 print(conf_matrix.rf)
56 cat('\n')
57
58 # matriz de confusão para o modelo SVM
59 conf_matrix.svm <- confusionMatrix(predict.svm, teste[["classes"]])
60
61 print(conf_matrix.svm)
62 cat('\n')
63
64 # matriz de confusão para o modelo RNA
65 conf_matrix.rna <- confusionMatrix(predict.rna, teste[["classes"]])
66
67 print(conf_matrix.rna)
68 cat('\n')
69
70 ##### 1.5 Indique qual modelo dá o melhor o resultado e a métrica
    utilizada
71
72 #O melhor modelo foi random forest com acurácia de 0.9182 e kappa de
    0.8987. A métrica utilizada foram a acurácia e kappa.

```

2 Estimativa de Volume de Árvores

```

1 ##### 2.1 Carregar o arquivo Volumes.csv (<http://www.razer.net.br/
    datasets/Volumes.csv>)
2
3 dados <- read.csv("http://www.razer.net.br/datasets/Volumes.csv", sep
    =";", dec=",")
4 head(dados)
5
6 ##### 2.2 Eliminar a coluna NR, que só apresenta um número sequencial
7
8 dados[["NR"]] <- NULL
9
10 ##### 2.3 Criar partição de dados: treinamento 80%, teste 20%

```

```

11
12 regression.indices <- caret::createDataPartition(dados[["VOL"]], p
    =0.8, list=F)
13 regression.treino <- dados[regression.indices, ]
14 regression.teste <- dados[-regression.indices, ]
15
16 ##### 2.4 Usando o pacote "caret", treinar os modelos: Random Forest (
    rf), SVM (svmRadial), Redes Neurais (neuralnet) e o modelo alomé
    trico de SPURR
17
18 # Para reproductibilidade
19 set.seed(7)
20
21 regression.rf <- caret::train(VOL ~ ., data=regression.treino, method
    ="rf")
22 regression.svm <- caret::train(VOL ~ ., data=regression.treino,
    method="svmRadial")
23 regression.rna <- caret::train(VOL ~ ., data=regression.treino,
    method="nnet", trControl=trainControl(method = "LOOCV"), trace=F)
24
25 ##### treino do modelo Spurr
26
27 regression.spurr <- nls( VOL ~ b0 + b1*DAP*DAP*HT, data=regression.
    treino, start=list(b0=0.5, b1=0.5))
28
29 ##### visualizar os resultados de Spurr
30
31 summary(regression.spurr)
32
33 ##### 2.5 Efetue as predições nos dados de teste
34
35 # predições
36
37 # Para reproductibilidade
38 set.seed(7)
39
40 predict.regression.rf <- predict(regression.rf, regression.teste)
41 predict.regression.svm <- predict(regression.svm, regression.teste)
42 predict.regression.rna <- predict(regression.rna, regression.teste)
43 predict.regression.suprr <- predict(regression.spurr, regression.
    teste)
44
45 ##### 2.6 Crie suas próprias funções (UDF) e calcule as seguintes mé
    tricas entre a predição e os dados observados
46
47 # - Erro padrão de estimativa: Syx
48
49 Syx <- function(reals, predicteds, n) {
50   return (sqrt(sum((reals - predicteds)^2)/(n - 2)))
51 }
52

```

```

53 # - Erro padrão de estimativa em porcentagem: Syx%
54
55 SyxPercent <- function(reals, predicted, n) {
56   return ((Syx(reals, predicted, n)/mean(reals))*100)
57 }
58
59 # - Coeficiente de determinação: R2
60
61 R2 <- function (reals, predicted) {
62   return (1 - sum((reals - predicted)^2)/sum((reals - mean(reals))^2))
63 }
64
65 ##### 2.7 Escolha o melhor modelo.
66
67 ##### métrica de estimativas para o modelo RandomForest - Regressão
68
69 # - coeficiente de determinação
70
71 R2(regression.teste[["VOL"]], predict.regression.rf)
72
73 # - Erro padrão estimativa
74
75 n <- nrow(regression.teste)
76 Syx(regression.teste[["VOL"]], predict.regression.rf, n)
77
78 # - Erro padrão estimativa em porcentagem
79
80 n <- nrow(regression.teste)
81 SyxPercent(regression.teste[["VOL"]], predict.regression.rf, n)
82
83 ##### métrica de estimativas para o modelo SVM - Regressão
84
85 # - coeficiente de determinação
86
87 R2(regression.teste[["VOL"]], predict.regression.svm)
88
89 # - Erro padrão estimativa
90
91 n <- nrow(regression.teste)
92 Syx(regression.teste[["VOL"]], predict.regression.svm, n)
93
94 # - Erro padrão estimativa em porcentagem
95
96 n <- nrow(regression.teste)
97 SyxPercent(regression.teste[["VOL"]], predict.regression.svm, n)
98
99 ##### métricas de estimativas para o modelo nnet - Regressão
100
101 # - coeficiente de determinação
102

```

```

103 R2(regression.teste[["VOL"]], predict.regression.rna)
104
105 # - Erro padrão estimativa
106
107 n <- nrow(regression.teste)
108 Syx(regression.teste[["VOL"]], predict.regression.rna, n)
109
110 # - Erro padrão estimativa em porcentagem
111
112 n <- nrow(regression.teste)
113 SyxPercent(regression.teste[["VOL"]], predict.regression.rna, n)
114
115 ##### métricas de estimativas para o modelo Spurr
116
117 # - coeficiente de determinação
118
119 R2(regression.teste[["VOL"]], predict.regression.suprr)
120
121 # - Erro padrão estimativa
122
123 n <- nrow(regression.teste)
124 Syx(regression.teste[["VOL"]], predict.regression.suprr, n)
125
126 # - Erro padrão estimativa em porcentagem
127
128 n <- nrow(regression.teste)
129 SyxPercent(regression.teste[["VOL"]], predict.regression.suprr, n)
130
131 ##### 2.7 escolha o melhor modelo
132
133 ##### Resumo dos resultados
134
135 ##### RF:
136
137 # 1. coeficiente de determinação: 0.8223603.
138 # 2. Erro padrão estimativa: 0.1376052.
139 # 3. Erro padrão estimativa em porcentagem: 10.42195
140
141 ##### SVM:
142
143 # 1. coeficiente de determinação: 0.6254546
144 # 2. Erro padrão estimativa: 0.19981
145 # 3. Erro padrão estimativa em porcentagem: 15.13322
146
147 ##### nnet:
148
149 # 1. coeficiente de determinação: -1.069672
150 # 2. Erro padrão estimativa: 0.4696948
151 # 3. Erro padrão estimativa em porcentagem: 35.57377
152
153 ##### spurr:

```

```
154 |  
155 | # 1. coeficiente de determinação: 0.7734018  
156 | # 2. Erro padrão estimativa: 0.1554151  
157 | # 3. Error padrão estimativa em porcentagem: 11.77084  
158 |  
159 | # Com base nas métricas, o modelo que se saiu melhor foi o  
    | RandomForest, com R2 igual 0.8223603, Erro padrão estimativa de  
    | 0.1376052 e Erro padrão de estima- tiva em porcentagem de 10.42195
```

APÊNDICE 4 – ESTATÍSTICA APLICADA I

A - ENUNCIADO

1. Gráficos e tabelas

(15 pontos) Elaborar os gráficos box-plot e histograma das variáveis “age” (idade da esposa) e “husage” (idade do marido) e comparar os resultados

(15 pontos) Elaborar a tabela de frequências das variáveis “age” (idade da esposa) e “husage” (idade do marido) e comparar os resultados

2. Medidas de posição e dispersão

(15 pontos) Calcular a média, mediana e moda das variáveis “age” (idade da esposa) e “husage” (idade do marido) e comparar os resultados

(15 pontos) Calcular a variância, desvio padrão e coeficiente de variação das variáveis “age” (idade da esposa) e “husage” (idade do marido) e comparar os resultados

3. Testes paramétricos ou não paramétricos

(40 pontos) Testar se as médias (se você escolher o teste paramétrico) ou as medianas (se você escolher o teste não paramétrico) das variáveis “age” (idade da esposa) e “husage” (idade do marido) são iguais, construir os intervalos de confiança e comparar os resultados.

Obs: Você deve fazer os testes necessários (e mostra-los no documento pdf) para saber se você deve usar o unpaired test (paramétrico) ou o teste U de Mann-Whitney (não paramétrico), justifique sua resposta sobre a escolha. Lembre-se de que os intervalos de confiança já são mostrados nos resultados dos testes citados no item 1 acima.

B - RESOLUÇÃO

Código

```
1 install.packages("BSDA")
2 install.packages("onewaytests")
3 install.packages("sjPlot")
4 install.packages("tigerstats")
5 install.packages("misty")
6 install.packages("dplyr")
7 install.packages("ggpubr")
8 install.packages("tidyverse")
9 install.packages("ggplot2")
10 install.packages("stats")
11
12 library("BSDA")
13 library("onewaytests")
14 library("sjPlot")
15 library("tigerstats")
16 library("misty")
17 library("dplyr")
```

```

18 library("ggpubr")
19 library("ggplot2")
20 library("tidyverse")
21 library("car")
22 library("fdth")
23 library("stats")
24 load("salarios.RData")
25
26 # 1 Gráficos e Tabelas
27
28 # a) Elaborar gráficos box-plot e histograma das variáveis "age"
29 (idade da esposa) e "husage" (idade do marido) e comparar os
    resultados
30 # Esposas
31 Boxplot( ~ age, data=salarios, id=list(method="y"), ylab="Distribuição
    o das Idade das Esposas")
32 # Conclusões: O gráfico demonstra a distribuição das idades das
    esposas
33 presentes no dataset salarios.
34 # Desta forma é possível perceber que há maior densidade da distribuiç
    ão
35 das idades entre os valores aproximados de 30 a 50 anos.
36 histograma.esposa <- hist(
37 salarios[["age"]],
38 xlab="Idade da Esposa",
39 ylab = "Frequência",
40 main = "Distribuição das idades das esposas"
41 )
42 # Conclusões: Quando observa-se a distribuição das frequências por
    meio
43 do gráfico de histograma para a idade das esposas,
44 # é possível notar a densidade observada no gráfico anterior entre os
45 intervalos de 30 a 50, no entanto, esse tipo de
46 # visualização fornece mais detalhamento quanto à presença de esposas
47 mais jovens e mais velhas.
48 # Maridos
49 Boxplot( ~ husage, data=salarios, id=list(method="y"), ylab="
    Distribuição
50 das Idade dos Maridos")
51 # Conclusões: A distribuição de idades demonstra a prevalência de
    homens
52 casados entre intervalos aproximados de 30 a 50 anos.
53 histograma.marido <- hist(
54 salarios[["husage"]],
55 xlab="Idade do Marido",
56 ylab = "Frequência",
57 main = "Distribuição das Idade dos Maridos"
58 )
59 # Conclusões: A distribuição das frequências das idades demonstram a
60 prevalência de maridos entre 30 a 50 anos.
61 # b) Elaborar a tabela de frequências das variáveis "age" (idade da

```

```

62 esposa) e "husage" (idade do marido) e comparar os resultados
63 freq.esposa <- table(salarios[["age"]])
64 # Conclusões: Observando a distribuição das frequências por idades
    das
65 esposas é possível observar que nesta base de dados as esposas
66 # mais novas possuem 18 anos e as mais velhas 59 anos. Outro fato
    curioso, é que conforme que as ocorrências de esposas mais novas são
67 # mais raras comparadas as mais velhas.
68 freq.marido <- table(salarios[["husage"]])
69 # Conclusões: Já o cenário dos homens demonstra que o mais novo
    possui 19
70 anos e o mais velho 86 anos. Outro fato é a pouca incidência
71 # de homens mais velhos na base de dados. E que comparativamente a
72 frequência de mulheres mais novas em relação a idade dos homens.
73
74 # 2 Medidas de Posição e Dispersão
75
76 # a) Calcular a média, mediana e moda das variáveis "age" (idade da
77 esposa) e "husage" (idade do marido) e comparar os resultados
78 # Esposa
79 media.esposa <- mean(salarios[["age"]])
80 media.esposa # 39.42 anos
81 mediana.esposa <- median(salarios[["age"]])
82 mediana.esposa # 39 anos
83 table(salarios[["age"]])
84 subset(table(salarios[["age"]]),
85 table(salarios[["age"]] == max(table(salarios[["age"]])))) # 217
    esposas com
86 37 anos
87 # Marido
88 media.marido <- mean(salarios[["husage"]])
89 media.marido # 42.45 anos
90 mediana.marido <- median(salarios[["husage"]])
91 mediana.marido # 41 anos
92 table(salarios[["husage"]])
93 subset(table(salarios[["husage"]]),
94 table(salarios[["husage"]] == max(table(salarios[["husage"]])))) #
    201
95 maridos com 44 anos
96
97 # Conclusões
98 media.marido/media.esposa # a média das idades maridos é 7,67%
    superior a
99 média das idades das esposas.
100 mediana.marido/mediana.esposa # a mediana das idades dos maridos é de
101 5,12% superior a mediana das idades das esposas.
102 44/37 # a moda das idades dos maridos é maior que a moda da idade das
103 esposas em 18,91%.
104 # b) Calcular a variância, desvio padrão e coeficiente de variação
    das

```

```

105 variáveis "age" (idade da esposa) e "husage" (idade do marido) e
    comparar
106 os resultados
107 # Esposa
108 variancia.esposa <- var(salarios[["age"]])
109 variancia.esposa # 99.75
110 desvio.esposa <- sd(salarios[["age"]])
111 desvio.esposa # +- 9.98 anos
112 coeV.esposa <- (desvio.esposa/media.esposa)*100
113 coeV.esposa # 25.33 %
114 # Marido
115 variancia.marido <- var(salarios[["husage"]])
116 variancia.marido # 126.07
117 desvio.marido <- sd(salarios[["husage"]])
118 desvio.marido # +- 11.22 anos
119 coeV.marido <- (desvio.marido/media.marido)*100
120 coeV.marido # 26.44 %
121 # Conclusões
122 variancia.marido/variancia.esposa # a variância das idades dos
    maridos é
123 26,38% superior ao das esposas.
124 desvio.marido/desvio.esposa # o desvio padrão das idades dos maridos
    é de
125 12,42% superior ao das esposas.
126 coeV.marido/coeV.esposa # o coeficiente de variação das idades dos
    maridos
127 é de 4,40% superior ao das esposas.
128
129 # 3 Testes paramétricos ou não paramétricos
130
131 # a) Testar se as médias (se você escolher o teste paramétrico) ou as
132 medianas (se você escolher o teste não paramétrico) das variáveis
133 # "age" (idade da esposa) e "husage" (idade do marido) são iguais,
134 construir os intervalos de confiança e comparar os resultados
135 marido.idades <- salarios[["husage"]]
136 esposa.idades <- salarios[["age"]]
137 group.marido <- rep(c('husband'), times=length(marido.idades))
138 group.esposa <- rep(c('wife'), times=length(esposa.idades))
139 data.marido <- data.frame(group=group.marido, ages=marido.idades)
140 data.esposa <- data.frame(group=group.esposa, ages=esposa.idades)
141 data.ages <- rbind(data.marido, data.esposa)
142 # Checando as informações sobre o dataframe gerado
143 group_by(data.ages, group) %>%
144 summarise(
145   count = n(),
146   mean = mean(ages, na.rm = TRUE),
147   median = median(ages, na.rm = TRUE),
148   IQR = IQR(ages, na.rm = TRUE)
149 )
150 ggboxplot(
151   data.ages,

```

```

152 x = 'group',
153 y = 'ages',
154 color = 'group',
155 palette=c("#00AFBB", "#E7B800"),
156 ylab = "Ages",
157 xlab = "Groups"
158 )
159 # Premissa 1: Os grupos são independentes? Sim, pois os grupos de
    idades
160 de maridos e esposas não estão relacionados. Assim não se encontram na
161 classificação de emparelhados.
162 # Premissa 2: Os dados possuem distribuição normal?
163 # H0: dados normalmente distribuídos
164 # H1: dados não normalmente distribuídos
165 options(scipen = 999)
166 amostra <- data.ages[data.ages[["group"]] == 'husband', ]
167 with(amostra, ks.test(unique(ages), "pnorm", mean = mean(ages), sd =
168 sd(ages)))
169 # p-value < 0.05 com rejeição de H0, dados não estão normalmente
170 distribuídos para idades dos maridos
171 amostra <- data.ages[data.ages[["group"]] == 'wife', ]
172 with(amostra, ks.test(unique(ages), "pnorm", mean = mean(ages), sd =
173 sd(ages)))
174 # p-value > 0.05 com aceite de H0, dados estão normalmente distribuí
    dos
175 para idades das esposas
176 # Premissa 3: As duas amostras possuem homogeneidade das variâncias?
177 # H0: as variâncias são estatisticamente iguais
178 # H1: as variâncias não são estatisticamente iguais
179 res.ftest <- var.test(ages ~ group, data = data.ages)
180 res.ftest
181 qf(0.95, 5633, 5633)
182 dist_f(f = 1.04481, deg.f1 = 5633, deg.f2 = 5633)
183 dist_f(f = 0.9571118, deg.f1 = 5633, deg.f2 = 5633)
184 #O teste de F tem valor critico entre 0.9571118 e 1.04481 (regiao de
    não
185 rejeição de H0),
186 #os valores acima de 1.04481 e abaixo de 0.9571118 estão na região de
187 rejeição de H0 (área azul do gráfico).
188 #O valor da estatística F calculada é 1.2638. Como esse valor se
    encontra
189 na região de rejeição de H0,
190 #então rejeitamos a hipótese de que as variâncias são
    estatisticamente
191 iguais.
192 # Utilizaremos então o teste não paramétrico de Mann-Whitney para
193 comparar as medianas entre os grupos.
194 # H0: a idade mediana dos maridos é estatisticamente igual a das
    esposas.
195
196 # H1: a idade mediana dos maridos não é estatisticamente igual a das

```

```

197 esposas.
198 res.w <- wilcox.test(ages ~ group, data = data.ages, exact = FALSE,
199 conf.int = TRUE)
200 res.w
201 # p-value < 0.05 com rejeição de H0, concluindo que a idade mediana
    dos
202 maridos não é estatisticamente igual à das esposas.
203 res.w <- wilcox.test(ages ~ group, data = data.ages, exact = FALSE,
204 conf.int = TRUE)
205 res.w
206 # Teste para verificar se a mediana da idade dos maridos é menor que a
207 mediana da idade das esposas.
208 # H0: a idade mediana dos maridos não é estatisticamente menor que a
209 idade mediana das esposas.
210 # H1: a idade mediana dos maridos é estatisticamente menor que a idade
211 mediana das esposas.
212 res.w <- wilcox.test(ages ~ group, data = data.ages, exact = FALSE,
213 conf.int = TRUE, alternative='less')
214 res.w
215 # p-value > 0.05, concluindo que a idade mediana dos maridos não é
216 estatisticamente menor que a idade mediana das esposas.
217 # Teste para verificar se a mediana da idade dos maridos é maior que a
218 mediana da idade das esposas.
219 # H0: a idade mediana dos maridos não é estatisticamente maior que a
220 idade mediana das esposas.
221 # H1: a idade mediana dos maridos é estatisticamente maior que a idade
222 mediana das esposas.
223 res.w <- wilcox.test(ages ~ group, data = data.ages, exact = FALSE,
224 conf.int = TRUE, alternative='greater')
225 res.w
226 # p-value < 0.05, concluindo que a idade mediana dos maridos não é
227 estatisticamente maior que a idade mediana das esposas.

```

APÊNDICE 5 – ESTATÍSTICA APLICADA II

A - ENUNCIADO

Regressões Ridge, Lasso e ElasticNet

(100 pontos) Fazer as regressões Ridge, Lasso e ElasticNet com a variável dependente “lwage” (salário-hora da esposa em logaritmo neperiano) e todas as demais variáveis da base de dados são variáveis explicativas (todas essas variáveis tentam explicar o salário-hora da esposa). No pdf você deve colocar a rotina utilizada, mostrar em uma tabela as estatísticas dos modelos (RMSE e R^2) e concluir qual o melhor modelo entre os três, e mostrar o resultado da predição com intervalos de confiança para os seguintes valores:

- husage = 40 (anos – idade do marido)
- husunion = 0 (marido não possui união estável)
- husearns = 600 (US\$ renda do marido por semana)
- huseduc = 13 (anos de estudo do marido)
- husblck = 1 (o marido é preto)
- hushisp = 0 (o marido não é hispânico)
- hushrs = 40 (horas semanais de trabalho do marido)
- kidge6 = 1 (possui filhos maiores de 6 anos)
- age = 38 (anos – idade da esposa)
- black = 0 (a esposa não é preta)
- educ = 13 (anos de estudo da esposa)
- hispanic = 1 (a esposa é hispânica)
- union = 0 (esposa não possui união estável)
- exper = 18 (anos de experiência de trabalho da esposa)
- kidlt6 = 1 (possui filhos menores de 6 anos)

Obs: lembre-se de que a variável dependente “lwage” já está em logaritmo, portanto você não precisa aplicar o logaritmo nela para fazer as regressões, mas é necessário aplicar o antilog para obter o resultado da predição.

B - RESOLUÇÃO 1 Regressões Ridge, Lasso e ElasticNet

A tabela a seguir apresenta os resultados obtidos no R ao processar os modelos Ridge, Lasso e ElasticNet.

Tabela 5.1: Desempenho dos modelos de regressão nas etapas de treino e teste.

Modelo	Etapa	RMSE	R^2	Lwage (US\$)	Lwage inferior (US\$)	Lwage superior (US\$)
Ridge	Treino	0.83	0.29	7.86	7.69	8.04
	Teste	0.92	0.23	7.86	7.69	8.04
Lasso	Treino	0.83	0.29	8.04	7.86	8.22
	Teste	0.93	0.23	8.04	7.86	8.22
ElasticNet	Treino	0.83	0.29	8.11	7.93	8.29
	Teste	0.92	0.22	8.11	7.93	8.29

Todos os modelos apresentaram $RMSE = 0,83$ e $R^2 = 0,29$, ou seja, o mesmo desempenho na etapa de treinamento. Já na etapa de teste, os modelos Ridge e ElasticNet ($RMSE = 0,92$) apresentaram desempenho preditivo ligeiramente superior ao do Lasso ($RMSE = 0,93$). O R^2 do ElasticNet foi de 0,22, valor ainda muito próximo ao dos demais modelos. O valor estimado do salário-hora da esposa foi de US\$ 7,86 para o Ridge, US\$ 8,04 para o Lasso e US\$ 8,11 para o ElasticNet. O modelo que apresentou o maior valor estimado foi o ElasticNet. As amplitudes dos intervalos de confiança foram semelhantes, variando aproximadamente entre US\$ 0,35 e US\$ 0,36, o que indica consistência entre os modelos.

Código

```

1 #####
2 # MODELO RIDGE
3 # Instalando os pacotes necessarios
4 install.packages("plyr")
5 install.packages("readr")
6 install.packages("dplyr")
7 install.packages("caret")
8 install.packages("ggplot2")
9 install.packages("repr")
10 install.packages("glmnet")
11 # Carregando os pacotes necessarios
12 library(plyr)
13 library(readr)
14 library(dplyr)
15 library(caret)
16 library(ggplot2)
17 library(repr)
18 library(glmnet)
19 # Vamos utilizar a base de dados "trabalhosalarios" que eh um dataset
20 # que contem uma amostra com dados sobre a renda de familias
21 # Carregando o dataset
22 load("~/Documents/ESPECIALIZACAO UFPR/ESTATISTICA 2/Bases de Dados
    Usadas nas Aulas Práticas/trabalhosalarios.RData" )
23 # Vamos guardar este dataset em um objeto chamaddo "dat"
24 dat <- trabalhosalarios
25 # Vamos ver o dataset inteiro
26 View(dat)
27 # Vamos visualizar parte do dataset
28 glimpse(dat)

```

```

29 # Vamos ver como esta nossa memoria
30 gc()
31 # Vamos jogar a semente para gerar numeros aleatorios
32 # Aqui no exemplo a semente eh "1106", mas poderia
33 # ser qualquer outro numero, se todos usarem a mesma
34 # semente os resultados serao iguais.
35 # Essa semente de numeros aleatorios serve para
36 # particionar o dataset aleatoriamente
37 set.seed(1106)
38 # Vamos criar um indice para particionar o dataset em
39 # 80% para treinamento
40 index = sample(1:nrow(dat), 0.8*nrow(dat))
41 # Vamos criar a base de dados de treinamento
42 train = dat[index,]
43 # Vamos criar a base de dados de teste
44 test = dat[-index,]
45 # Vamos checar as dimensoes das bases de treinamento e
46 # teste
47 dim(train)
48 dim(test)
49 # A base de treinamento possui 2059 linhas (familias) e
50 # 17 colunas (variaveis)
51 # A base de teste possui 515 linhas (familias) e 17
52 # colunas (variaveis)
53 # Vamos padronizar as variaveis
54 # Vamos criar um objeto com as variaveis para padronizar
55 # As variaveis binarias nao sao padronizadas
56 cols = c('husage', 'husearns', 'huseduc', 'hushrs', 'earns',
57 'age', 'educ', 'exper', 'lwage')
58 colnames(train)
59 # Padronizando a base de treinamento e teste
60 pre_proc_val <- preProcess(train[,cols], method = c("center", "scale"
61 ))
62 train[,cols] = predict(pre_proc_val, train[,cols])
63 test[,cols] = predict(pre_proc_val, test[,cols])
64 # Vamos ver o sumario estatistico das variaveis
65 # padronizadas de cada dataset
66 summary(train)
67 summary(test)
68 #####
69 # REGRESSAO RIDGE #
70 #####
71 # A regressao Ridge "encolhe os valores dos coeficientes"
72 # Vamos criar um objeto com as variaveis que usaremos no
73 # modelo
74 cols_reg = c('husage',
75 'husearns',
76 'huseduc',
77 'hushrs',
78 'age',

```

```

79     'educ',
80     'husblck',
81     'hushisp',
82     'kidge6',
83     'black',
84     'hispanic',
85     'union',
86     'kidlt6',
87     'exper',
88     'husunion',
89     'lwage')
90 # Vamos gerar variaveis dummies para organizar os datasets
91 # em objetos tipo matriz
92 # Estamos interessados em estimar o salario-hora (lwage)
93 dummies <- dummyVars(lwage~usage+
94     husunion+
95     husearns+
96     huseduc+
97     hushrs+
98     husblck+
99     hushisp+
100     kidge6+
101     age+
102     educ+
103     black+
104     hispanic+
105     union+
106     kidlt6+
107     exper,
108     data = dat[,cols_reg])
109 train_dummies = predict(dummies, newdata = train[,cols_reg])
110 test_dummies = predict(dummies, newdata = test[,cols_reg])
111 print(dim(train_dummies)); print(dim(test_dummies))
112 # A regressao Ridge eh uma extensao da regressao linear em
113 # que a funcao perda eh alterada para minimizar a
114 # complexidade do modelo.
115 # Esta mudanca eh feita ao adicionarmos um parametro de
116 # penalty igual ao quadrado dos valores dos coeficientes.
117 # A principal diferenca entre a regressao linear e a
118 # penalizada eh que a regressao penalizada usa um
119 # hiperparametro "lambda".
120 # A funcao glmnet() executa o modelo muitas vezes para
121 # valores diferentes de lambda.
122 # Podemos automatizar esta tarefa para encontrar o melhor
123 # valor de lambda usando a funcao "cv.glmnet()".
124 # A funcao de perda generica eh:
125 # Funcao perda = MQO+lambda*soma(quadrados dos valores
126 # dos coeficientes).
127 # O lambda eh o parametro de penalty encontrado.
128 # Vamos guardar a matriz de dados de treinamento das
129 # variaveis explicativas para o modelo em um objeto

```

```

130 # chamado "x"
131 x = as.matrix(train_dummies)
132 # Vamos guardar o vetor de dados de treinamento da
133 # variavel dependente para o modelo em um objeto
134 # chamado "y_train"
135 y_train = train[["lwage"]]
136 # Vamos guardar a matriz de dados de teste das variaveis
137 # explicativas para o modelo em um objeto chamado
138 # "x_test"
139 x_test = as.matrix(test_dummies)
140 # Vamos guardar o vetor de dados de teste da variavel
141 # dependente para o modelo em um objeto chamado "y_test"
142 y_test = test[["lwage"]]
143 # Vamos calcular o valor otimo de lambda;
144 # alpha = "0", eh para regressao Ridge
145 # Vamos testar os lambdas de 10^-3 ate 10^2, a cada 0.1
146 lambdas <- 10^seq(2, -3, by = -.1)
147 # Calculando o lambda:
148 ridge_lamb <- cv.glmnet(x, y_train, alpha = 0,
149 lambda = lambdas)
150 # Vamos ver qual o lambda otimo
151 best_lambda_ridge <- ridge_lamb[["lambda"]].min
152 best_lambda_ridge
153 # Para contar o tempo de processamento do modelo usamos:
154 start <- Sys.time()
155 # Estimando o modelo Ridge
156 ridge_reg = glmnet(x, y_train, nlambda = 25, alpha = 0,
157 family = 'gaussian',
158 lambda = best_lambda_ridge)
159 end <- Sys.time()
160 difftime(end, start, units="secs")
161 # Vamos ver o resultado (valores) da estimativa
162 # (coeficientes)
163 ridge_reg[["beta"]]
164 # Vamos calcular o R^2 dos valores verdadeiros e
165 # preditos conforme a seguinte funcao:
166 eval_results <- function(true, predicted, df) {
167   SSE <- sum((predicted - true)^2)
168   SST <- sum((true - mean(true))^2)
169   R_square <- 1 - SSE / SST
170   RMSE = sqrt(SSE/nrow(df))
171   # As metricas de performace do modelo:
172   data.frame(
173     RMSE = RMSE,
174     Rsquare = R_square
175   )
176 }
177 # Predicao e avaliacao nos dados de treinamento:
178 predictions_train <- predict(ridge_reg,
179 s = best_lambda_ridge,
180 newx = x)

```

```

181 # As metricas da base de treinamento sao:
182 eval_results(y_train, predictions_train, train)
183 # Predicao e avaliacao nos dados de teste:
184 predictions_test <- predict(ridge_reg,
185 s = best_lambda_ridge,
186 newx = x_test)
187 # As metricas da base de teste sao:
188 eval_results(y_test, predictions_test, test)
189 # Vamos fazer uma predicao para o nosso exemplo
190 # Como os valores do dataset sao padronizados, nos temos
191 # de padronizar tambem os dados que vamos fazer a predicao
192 # Note que as variaveis dummies nao devem ser padronizadas
193 # Vamos fazer uma predicao para:
194 # husage = 40 anos (idade do marido)
195 husage = (40-pre_proc_val[["mean"]][["husage"]])/pre_proc_val[["std"
  ]][["husage"]]
196 husunion = 0
197 # husearns = 551 (rendimento do marido em dollar)
198 husearns = (600-pre_proc_val[["mean"]][["husearns"]])/pre_proc_val[["
  std"]]
199 [["husearns"]]
200 # huseduc = 13 (anos de estudo do marido)
201 huseduc = (13-pre_proc_val[["mean"]][["huseduc"]])/pre_proc_val[["std
  "]][["huseduc"]]
202 # husblk = 1 (o marido eh preto)
203 husblk = 1
204 # hushisp = 0 (o marido nao eh hispanico)
205 hushisp = 0
206 # hushrs = 40 (o marido trabalha 40 horas semanais)
207 hushrs = (40-pre_proc_val[["mean"]][["hushrs"]])/pre_proc_val[["std"
  ]][["hushrs"]]
208 # kidge6 = 1 (tem filhos maiores de 6 anos)
209 kidge6 = 1
210 # age = 38 anos (idade da esposa)
211 age = (38-pre_proc_val[["mean"]][["age"]])/pre_proc_val[["std"]][["
  age"]]
212 # black = 0 (esposa nao eh preta)
213 black = 0
214 # educ = 13 (esposa possui 13 anos de estudo)
215 educ = (13-pre_proc_val[["mean"]][["educ"]])/pre_proc_val[["std"]][["
  educ"]]
216 # hispanic = 0 (esposa nao eh hispanica)
217 hispanic = 0
218 # union = 0 (o casal nao possui uniao registrada)
219 union = 0
220 # exper = 18 anos de experiencia de trabalho da esposa
221 exper = (18-pre_proc_val[["mean"]][["exper"]])/pre_proc_val[["std"
  ]][["exper"]]
222 # kidlt6 = 1 (possui filhos com menos de 6 anos)
223 kidlt6 = 1
224 dataframe <- data.frame(husage=husage,

```

```

225     husunion=husunion,
226     husearns=husearns,
227     huseduc=huseduc,
228     husblk=husblk,
229     hushisp=hushisp,
230     hushrs=hushrs,
231     kidge6=kidge6,
232     exper=exper,
233     age=age,
234     black=black,
235     educ=educ,
236     hispanic=hispanic,
237     union=union,
238     kidlt6=kidlt6)
239 View(dataframe)
240 colnames(dataframe)
241 # Vamos construir uma matriz de dados para a predicao
242 our_pred = as.matrix(dataframe)
243 # Fazendo a predicao:
244 predict_our_ridge <- predict(ridge_reg,
245 s = best_lambda_ridge,
246 newx = our_pred)
247 # O resultado da predicao eh:
248 predict_our_ridge
249 # O resultado eh um valor padronizado, vamos converte-lo
250 # para o valor nominal, consistente com o dataset original
251 wage_pred_ridge=(predict_our_ridge*
252 pre_proc_val[["std"]][["lwage"]])+
253 pre_proc_val[["mean"]][["lwage"]])
254 # O resultado eh:
255 exp(wage_pred_ridge)
256 # Este eh o valor predito do salario por hora (US$),
257 # segundo as caracteristicas que atribuimos
258 # O intervalo de confianca para o nosso exemplo eh:
259 n <- nrow(train) # tamanho da amostra
260 m <- wage_pred_ridge # valor medio predito
261 s <- pre_proc_val[["std"]][["lwage"]] # desvio padrao
262 dam <- s/sqrt(n) # distribuicao da amostragem da media
263 CIlwr_ridge <- m + (qnorm(0.025))*dam # intervalo inferior
264 CIupr_ridge <- m - (qnorm(0.025))*dam # intervalo superior
265 # Os valores sao:
266 exp(CIlwr_ridge)
267 exp(CIupr_ridge)
268 # Entao, segundo as caracteristicas que atribuimos o
269 # salario-hora da esposa eh em media US$7.86 e pode
270 # variar entre US$7.69 e US$8.04

```

```

1 #####
2 # MODELO LASSO
3 # Carregando os pacotes necessarios
4 library(plyr)

```

```

5 library(readr)
6 library(dplyr)
7 library(caret)
8 library(ggplot2)
9 library(repr)
10 library(glmnet)
11 # Carregando o dataset
12 load("~/Documents/ESPECIALIZACAO UFPR/ESTATISTICA 2/Bases de Dados
    Usadas nas Aulas Práticas/trabalhosalarios.RData" )
13 # Vamos guardar este dataset em um objeto chamaddo "dat"
14 dat <- trabalhosalarios
15 # Vamos ver o dataset inteiro
16 View(dat)
17 # Vamos visualizar parte do dataset
18 glimpse(dat)
19 # Vamos ver como esta nossa memoria
20 gc()
21 # Vamos jogar a semente para gerar numeros aleatorios
22 # Aqui no exemplo a semente eh "1106", mas poderia
23 # ser qualquer outro numero, se todos usarem a mesma
24 # semente os resultados serao iguais.
25 # Essa semente de numeros aleatorios serve para
26 # particionar o dataset aleatoriamente
27 set.seed(1106)
28 # Vamos criar um indice para particionar o dataset em
29 # 80% para treinamento
30 index = sample(1:nrow(dat),0.8*nrow(dat))
31 # Vamos criar a base de dados de treinamento
32 train = dat[index,]
33 # Vamos criar a base de dados de teste
34 test = dat[-index,]
35 # Vamos checar as dimensoes das bases de treinamento e
36 # teste
37 dim(train)
38 dim(test)
39 # Vamos padronizar as variaveis
40 # Vamos criar um objeto com as variaveis para padronizar
41 # As variaveis binarias nao sao padronizadas
42 cols = c('husage', 'husearns', 'huseduc', 'hushrs', 'earns',
43 'age', 'educ', 'exper', 'lwage')
44 # Padronizando a base de treinamento e teste
45 pre_proc_val <- preProcess(train[,cols],
46 method = c("center", "scale"))
47 train[,cols] = predict(pre_proc_val, train[,cols])
48 test[,cols] = predict(pre_proc_val, test[,cols])
49 # Vamos ver o sumario estatistico das variaveis
50 # padronizadas de cada dataset
51 summary(train)
52 summary(test)
53
54 #####

```

```

55 # REGRESSAO LASSO #
56 #####
57 # Vamos criar um objeto com as variaveis que usaremos no
58 # modelo
59 cols_reg = c('husage',
60 'husearns',
61 'huseduc',
62 'hushrs',
63 'age',
64 'educ',
65 'husblck',
66 'hushisp',
67 'kidge6',
68 'black',
69 'hispanic',
70 'union',
71 'kidlt6',
72 'exper',
73 'husunion',
74 'lwage')
75 # Vamos gerar variaveis dummies para organizar os datasets
76 # em objetos tipo matriz
77 # Estamos interessados em estimar o salario-hora (hrwage)
78 dummies <- dummyVars(lwage~husage+
79   husunion+
80   husearns+
81   huseduc+
82   hushrs+
83   husblck+
84   hushisp+
85   kidge6+
86   age+
87   educ+
88   black+
89   hispanic+
90   union+
91   kidlt6+
92   exper,
93   data = dat[,cols_reg])
94 train_dummies = predict(dummies, newdata = train[,cols_reg])
95 test_dummies = predict(dummies, newdata = test[,cols_reg])
96 print(dim(train_dummies)); print(dim(test_dummies))
97 # Vamos guardar a matriz de dados de treinamento das
98 # variaveis explicativas para o modelo em um objeto
99 # chamado "x"
100 x = as.matrix(train_dummies)
101 # Vamos guardar o vetor de dados de treinamento da
102 # variavel dependente para o modelo em um objeto
103 # chamado "y_train"
104 y_train = train[["lwage"]]
105 # Vamos guardar a matriz de dados de teste das variaveis

```

```

106 # explicativas para o modelo em um objeto chamado
107 # "x_test"
108 x_test = as.matrix(test_dummies)
109 # Vamos guardar o vetor de dados de teste da variavel
110 # dependente para o modelo em um objeto chamado "y_test"
111 y_test = test[["lwage"]]
112 # A regressao Lasso reduz os coeficientes nao significativos
113 # a "zero"
114 # A regressao Lasso, ou Operador de Encolhimento Absoluto
115 # Minimo e de Selecao, tambem eh uma variacao da regressao
116 # linear.
117 # Para a regressao Lasso, a funcao perda eh alterada para
118 # minimizar a complexidade do modelo, restringindo a soma
119 # dos valores absolutos dos coeficientes do modelo
120 # (tambem chamada de restricao "l1-norm").
121 # A restricao "l1-norm" forca alguns alguns valores dos
122 # pesos a "zero" para permitir que outros coeficientes
123 # tenham valores mais distantes de "zero".
124 # A funcao perda eh:
125 # Funcao perda = MQO+lambda*soma(valores absolutos dos
126 # coeficientes)
127 # Vamos criar o intervalo para o tuning do melhor lambda
128 lambdas <- 10^seq(2, -3, by = -.1)
129 # Vamos atribuir alpha = 1 para implementar a regressao
130 # lasso
131 lasso_lamb <- cv.glmnet(x, y_train, alpha = 1,
132 lambda = lambdas,
133 standardize = TRUE, nfolds = 5)
134 # Vamos guardar o lambda "otimo" em um objeto chamado
135 # best_lambda_lasso
136 best_lambda_lasso <- lasso_lamb[["lambda"]].min
137 best_lambda_lasso
138 # Vamos estimar o modelo Lasso
139 lasso_model <- glmnet(x, y_train, alpha = 1,
140 lambda = best_lambda_lasso,
141 standardize = TRUE)
142 # Vamos visualizar os coeficientes estimados
143 lasso_model[["beta"]]
144 # Perceba que alguns coeficientes estao zerados pois
145 # nao sao significativos
146 # Vamos fazer as predicoes na base de treinamento e
147 # avaliar a regressao Lasso
148 predictions_train <- predict(lasso_model,
149 s = best_lambda_lasso,
150 newx = x)
151 # Vamos calcular o R^2 dos valores verdadeiros e
152 # preditos conforme a seguinte funcao:
153 eval_results <- function(true, predicted, df) {
154   SSE <- sum((predicted - true)^2)
155   SST <- sum((true - mean(true))^2)
156   R_square <- 1 - SSE / SST

```

```

157 RMSE = sqrt(SSE/nrow(df))
158 # As metricas de performace do modelo:
159 data.frame(RMSE = RMSE, Rsquare = R_square)
160 }
161 # As metricas da base de treinamento sao:
162 eval_results(y_train, predictions_train, train)
163 # Vamos fazer as predicoes na base de teste
164 predictions_test <- predict(lasso_model,
165 s = best_lambda_lasso,
166 newx = x_test)
167 # As metricas da base de teste sao:
168 eval_results(y_test, predictions_test, test)
169 # Perceba que os erros nao sao muito diferentes, assim
170 # como os valores de R2.
171 # Vamos fazer uma predicao para o nosso exemplo
172 # Vamos fazer a predicao com base nos mesmos parametros
173 # como fizemos anteriormente na regressao Ridge
174 # Vamos fazer uma predicao para:
175 # husage = 40 anos (idade do marido)
176 husage = (40-pre_proc_val[["mean"]][["husage"]])/pre_proc_val[["std"
  ]][["husage"]]
177 husunion = 0
178 # husearns = 551 (rendimento do marido em USS)
179 husearns = (600-pre_proc_val[["mean"]][["husearns"]])/pre_proc_val[["
  std"]]
180 [["husearns"]]
181 # huseduc = 13 (anos de estudo do marido)
182 huseduc = (13-pre_proc_val[["mean"]][["huseduc"]])/pre_proc_val[["std
  "]][["huseduc"]]
183 # husblk = 1 (o marido eh preto)
184 husblk = 1
185 # hushisp = 0 (o marido nao eh hispanico)
186 hushisp = 0
187 # hushrs = 40 (o marido trabalha 40 horas semanais)
188 hushrs = (40-pre_proc_val[["mean"]][["hushrs"]])/pre_proc_val[["std"
  ]][["hushrs"]]
189 # kidge6 = 1 (tem filhos maiores de 6 anos)
190 kidge6 = 1
191 # age = 38 anos (idade da esposa)
192 age = (38-pre_proc_val[["mean"]][["age"]])/pre_proc_val[["std"]][["
  age"]]
193 # black = 0 (esposa nao eh preta)
194 black = 0
195 # educ = 13 (esposa possui 13 anos de estudo)
196 educ = (13-pre_proc_val[["mean"]][["educ"]])/pre_proc_val[["std"]][["
  educ"]]
197 # hispanic = 0 (esposa nao eh hispanica)
198 hispanic = 0
199 # union = 0 (o casal nao possui uniao registrada)
200 union = 0
201 # exper = 18 anos de experiencia de trabalho da esposa

```

```

202 exper = (18-pre_proc_val[["mean"]][["exper"]])/pre_proc_val[["std"
    ]][["exper"]]
203 # kidlt6 = 1 (possui filhos com menos de 6 anos)
204 kidlt6 = 1
205 dataframe <- data.frame(husage=husage,
206     husunion=husunion,
207     husearns=husearns,
208     huseduc=huseduc,
209     husblck=husblck,
210     hushisp=hushisp,
211     hushrs=hushrs,
212     kidge6=kidge6,
213     exper=exper,
214     age=age,
215     black=black,
216     educ=educ,
217     hispanic=hispanic,
218     union=union,
219     kidlt6=kidlt6)
220 View(dataframe)
221 colnames(dataframe)
222 # Vamos construir uma matriz de dados para a predicao
223 our_pred = as.matrix(dataframe)
224 # Vamos para a predicao
225 predict_our_lasso <- predict(lasso_model,
226     s = best_lambda_lasso,
227     newx = our_pred)
228 predict_our_lasso
229 # Novamente, o resultado esta padronizado, nos temos de
230 # converte-lo para valor compativel com o dataset original
231 wage_pred_lasso=(predict_our_lasso*
232     pre_proc_val[["std"]][["lwage"]])+
233     pre_proc_val[["mean"]][["lwage"]]
234 exp(wage_pred_lasso)
235 # Portanto, o salario-hora da esposa eh USS 2.09
236 # Vamos criar o intervalo de confianca para o nosso
237 # exemplo
238 n <- nrow(train)
239 m <- wage_pred_lasso
240 s <- pre_proc_val[["std"]][["lwage"]]
241 dam <- s/sqrt(n)
242 CIlwr_lasso <- m + (qnorm(0.025))*dam
243 CIupr_lasso <- m - (qnorm(0.025))*dam
244 # O intervalo de confianca eh:
245 exp(CIlwr_lasso)
246 exp(CIupr_lasso)
247 # Entao, o salario hora da esposa eh de USS 8.04 e pode variar
248 # entre USS 7.93 e USS 8.29

```

```

1 #####
2 # MODELO ELASTIC NET

```

```

3 # Carregando os pacotes necessarios
4 library(plyr)
5 library(readr)
6 library(dplyr)
7 library(ggplot2)
8 library(repr)
9 library(glmnet)
10 library(caret)
11 # Carregando o dataset
12 load("~/Documents/ESPECIALIZACAO UFPR/ESTATISTICA 2/Bases de Dados
    Usadas nas Aulas Práticas/trabalhosalarios.RData" )
13 # Vamos guardar este dataset em um objeto chamado "dat"
14 dat <- trabalhosalarios
15 # Vamos ver o dataset inteiro
16 View(dat)
17 # Vamos visualizar parte do dataset
18 glimpse(dat)
19 # Vamos ver como esta nossa memoria
20 gc()
21 # Vamos jogar a semente para gerar numeros aleatorios
22 # Aqui no exemplo a semente eh "1106", mas poderia
23 # ser qualquer outro numero, se todos usarem a mesma
24 # semente os resultados serao iguais.
25 # Essa semente de numeros aleatorios serve para
26 # particionar o dataset aleatoriamente
27 set.seed(1106)
28 # Vamos criar um indice para particionar o dataset em
29 # 80% para treinamento
30 index = sample(1:nrow(dat), 0.8*nrow(dat))
31 # Vamos criar a base de dados de treinamento
32 train = dat[index,]
33 # Vamos criar a base de dados de teste
34 test = dat[-index,]
35 # Vamos checar as dimensoes das bases de treinamento e
36 # teste
37 dim(train)
38 dim(test)
39 # Vamos padronizar as variaveis
40 # Vamos criar um objeto com as variaveis para padronizar
41 # As variaveis binarias nao sao padronizadas
42 cols = c('husage', 'husearns', 'huseduc', 'hushrs', 'earns',
43 'age', 'educ', 'exper', 'lwage')
44 colnames(train)
45 # Padronizando a base de treinamento e teste
46 pre_proc_val <- preProcess(train[,cols],
47 method = c("center", "scale"))
48 train[,cols] = predict(pre_proc_val, train[,cols])
49 test[,cols] = predict(pre_proc_val, test[,cols])
50 # Vamos ver o sumario estatistico das variaveis
51 # padronizadas de cada dataset
52 summary(train)

```

```

53 summary(test)
54
55 #####
56 # REGRESSAO ELASTICNET #
57 #####
58 # Vamos criar um objeto com as variaveis que usaremos no
59 # modelo
60 cols_reg = c('husage',
61             'husearns',
62             'huseduc',
63             'hushrs',
64             'age',
65             'educ',
66             'husblack',
67             'hushisp',
68             'kidge6',
69             'black',
70             'hispanic',
71             'union',
72             'kidlt6',
73             'exper',
74             'husunion',
75             'lwage')
76 # Vamos gerar variaveis dummies para organizar os datasets
77 # em objetos tipo matriz
78 # Estamos interessados em estimar o salario-hora (lwage)
79 dummies <- dummyVars(lwage~husage+
80                     husunion+
81                     husearns+
82                     huseduc+
83                     hushrs+
84                     husblack+
85                     hushisp+
86                     kidge6+
87                     age+
88                     educ+
89                     black+
90                     hispanic+
91                     union+
92                     kidlt6+
93                     exper,
94                     data = dat[,cols_reg])
95 train_dummies = predict(dummies, newdata = train[,cols_reg])
96 test_dummies = predict(dummies, newdata = test[,cols_reg])
97 print(dim(train_dummies)); print(dim(test_dummies))
98 # Vamos guardar a matriz de dados de treinamento das
99 # variaveis explicativas para o modelo em um objeto
100 # chamado "x"
101 x = as.matrix(train_dummies)
102 # Vamos guardar o vetor de dados de treinamento da
103 # variavel dependente para o modelo em um objeto

```

```

104 # chamado "y_train"
105 y_train = train[["lwage"]]
106 # Vamos guardar a matriz de dados de teste das variaveis
107 # explicativas para o modelo em um objeto chamado
108 # "x_test"
109 x_test = as.matrix(test_dummies)
110 # Vamos guardar o vetor de dados de teste da variavel
111 # dependente para o modelo em um objeto chamado "y_test"
112 y_test = test[["lwage"]]
113 # A regressao ElasticNet combina aspectos das regressoes
114 # Ridge e Lasso. Ela penaliza o modelo usando ambas
115 # funcoes de penalidade, l2-norm and l1-norm.
116 # O modelo pode ser customizado pelo pacote "caret"
117 # que encontra os valores otimos dos parametros
118 # automaticamente.
119 # Vamos configurar o treinamento do modelo por
120 # cross validation, com 10 folders, 5 repeticoes
121 # e busca aleatoria dos componentes das amostras
122 # de treinamento, o "verboseIter" eh soh para
123 # mostrar o processamento.
124 train_cont <- trainControl(method = "repeatedcv",
125 number = 10,
126 repeats = 5,
127 search = "random",
128 verboseIter = TRUE)
129 # Nos nao temos o parametro "alpha", porque a regressao
130 # ElasticNet vai encontra-lo automaticamente, cujo valor
131 # estara entre 0 e 1 (para Ridge ==> alpha = 0; e
132 # para Lasso ==> alpha = 1).
133 # O parametro "lambda" tambem eh escolhido por
134 # cross-validation
135 # Vamos treinar o modelo
136 elastic_reg <- train(lwage~husage+
137   husunion+
138   husearns+
139   huseduc+
140   hushrs+
141   husblck+
142   hushisp+
143   kidge6+
144   age+
145   educ+
146   black+
147   hispanic+
148   union+
149   kidlt6+
150   exper,
151   data = train,
152   method = "glmnet",
153   tuneLength = 10,
154   trControl = train_cont)

```

```

155 # O melhor parametro alpha escolhido eh:
156 elastic_reg[["bestTune"]]
157 # E os parametros sao:
158 elastic_reg[["finalModel"]][["beta"]]
159 # Vamos fazer as predicoes e avaliar a performance do
160 # modelo
161 # Vamos fazer as predicoes no modelo de treinamento:
162 predictions_train <- predict(elastic_reg, x)
163 # Vamos calcular o R^2 dos valores verdadeiros e
164 # preditos conforme a seguinte funcao:
165 eval_results <- function(true, predicted, df) {
166   SSE <- sum((predicted - true)^2)
167   SST <- sum((true - mean(true))^2)
168   R_square <- 1 - SSE / SST
169   RMSE = sqrt(SSE/nrow(df))
170   # As metricas de performace do modelo:
171   data.frame(RMSE = RMSE, Rsquare = R_square)
172 }
173 # As metricas de performance na base de treinamento
174 # sao:
175 eval_results(y_train, predictions_train, train)
176 # Vamos fazer as predicoes na base de teste
177 # Vamos fazer uma predicao para:
178 # husage = 40 anos (idade do marido)
179 husage = (40-pre_proc_val[["mean"]][["husage"]])/pre_proc_val[["std"
  ]][["husage"]]
180 husunion = 0
181 # husearns = 551 (rendimento do marido em USS)
182 husearns = (600-pre_proc_val[["mean"]][["husearns"]])/pre_proc_val[["
  std"]]
183 [["husearns"]]
184 # huseduc = 13 (anos de estudo do marido)
185 huseduc = (13-pre_proc_val[["mean"]][["huseduc"]])/pre_proc_val[["std
  "]][["huseduc"]]
186 # husblk = 1 (o marido eh preto)
187 husblk = 1
188 # hushisp = 0 (o marido nao eh hispanico)
189 hushisp = 0
190 # hushrs = 40 (o marido trabalha 40 horas semanais)
191 hushrs = (40-pre_proc_val[["mean"]][["hushrs"]])/pre_proc_val[["std"
  ]][["hushrs"]]
192 # kidge6 = 1 (tem filhos maiores de 6 anos)
193 kidge6 = 1
194 # age = 38 anos (idade da esposa)
195 age = (38-pre_proc_val[["mean"]][["age"]])/pre_proc_val[["std"]][["
  age"]]
196 # black = 0 (esposa nao eh preta)
197 black = 0
198 # educ = 13 (esposa possui 13 anos de estudo)
199 educ = (13-pre_proc_val[["mean"]][["educ"]])/pre_proc_val[["std"]][["
  educ"]]

```

```

200 # hispanic = 0 (esposa nao eh hispanica)
201 hispanic = 0
202 # union = 0 (o casal nao possui uniao registrada)
203 union = 0
204 # exper = 18 anos de experiencia de trabalho da esposa
205 exper = (18-pre_proc_val[["mean"]][["exper"]])/pre_proc_val[["std"
    ]][["exper"]]
206 # kidlt6 = 1 (possui filhos com menos de 6 anos)
207 kidlt6 = 1
208 dataframe <- data.frame(husage=husage,
209     husunion=husunion,
210     husearns=husearns,
211     huseduc=huseduc,
212     husblck=husblck,
213     hushisp=hushisp,
214     hushrs=hushrs,
215     kidge6=kidge6,
216     exper=exper,
217     age=age,
218     black=black,
219     educ=educ,
220     hispanic=hispanic,
221     union=union,
222     kidlt6=kidlt6)
223 View(dataframe)
224 colnames(dataframe)
225 # Vamos construir uma matriz de dados para a predicao
226 our_pred = as.matrix(dataframe)
227 # Vamos fazer a predicao com base nos parametros que
228 # selecionamos
229 predict_our_elastic <- predict(elastic_reg,our_pred)
230 predict_our_elastic
231 # Novamente, o resultado eh padronizado, nos temos que
232 # reverte-lo para o nivel dos valores originais do
233 # dataset, vamos fazer isso:
234 wage_pred_elastic=(predict_our_elastic*
235 pre_proc_val[["std"]][["lwage"]])+
236 pre_proc_val[["mean"]][["lwage"]]
237 exp(wage_pred_elastic)
238 # Vamos criar o intervalo de confianca para o nosso
239 # exemplo
240 n <- nrow(train)
241 m <- wage_pred_elastic
242 s <- pre_proc_val[["std"]][["lwage"]]
243 dam <- s/sqrt(n)
244 CIlwr_elastic <- m + (qnorm(0.025))*dam
245 CIupr_elastic <- m - (qnorm(0.025))*dam
246 # Os valores minimo e maximo sao:
247 exp(CIlwr_elastic)
248 exp(CIupr_elastic)
249 # Entao, o salario-hora medio da esposa eh de USS8.11

```

```
250 # e pode variar entre USS 7.93 e USS 8.29
251 #####
```

APÊNDICE 6 – ARQUITETURA DE DADOS

A - ENUNCIADO

1 Construção de Características: Identificador automático de idioma

O problema consiste em criar um modelo de reconhecimento de padrões que dado um texto de entrada, o programa consegue classificar o texto e indicar a língua em que o texto foi escrito.

Parta do exemplo (notebook produzido no Colab) que foi disponibilizado e crie as funções para calcular as diferentes características para o problema da identificação da língua do texto de entrada.

Nessa atividade é para "construir características".

Meta: a acurácia deverá ser maior ou igual a 70%.

Essa tarefa pode ser feita no Colab (Google) ou no Jupiter, em que deverá exportar o notebook e imprimir o notebook para o formato PDF. Envie no UFPR Virtual os dois arquivos.

2 Melhore uma base de dados ruim

Escolha uma base de dados pública para problemas de classificação, disponível ou com origem na UCI Machine Learning.

Use o mínimo de intervenção para rodar a SVM e obtenha a matriz de confusão dessa base.

O trabalho começa aqui, escolha as diferentes tarefas discutidas ao longo da disciplina, para melhorar essa base de dados, até que consiga efetivamente melhorar o resultado.

Considerando a acurácia para bases de dados balanceadas ou quase balanceadas, se o percentual da acurácia original estiver em até 85%, a meta será obter 5%. Para bases com mais de 90% de acurácia, a meta será obter a melhora em pelo menos 2 pontos percentuais (92% ou mais).

Nessa atividade deverá ser entregue o script aplicado (o notebook e o PDF correspondente).

B - RESOLUÇÃO

1 Construção de Características: Identificador automático de idioma

```
1 ingles = [  
2     "Hello, how are you?",  
3     "I love to read books.",  
4     "The weather is nice today.",  
5     "Where is the nearest restaurant?",  
6     "What time is it?",  
7     "I enjoy playing soccer.",  
8     "Can you help me with this?",  
9     "I'm going to the movies tonight.",  
10    "This is a beautiful place.",  
11    "I like listening to music.",  
12    "Do you speak English?",  
13    "What is your favorite color?",  
14    "I'm learning to play the guitar.",  
15    "Have a great day!",
```

```

16 "I need to buy some groceries.",
17 "Let's go for a walk.",
18 "How was your weekend?",
19 "I'm excited for the concert.",
20 "Could you pass me the salt, please?",
21 "I have a meeting at 2 PM.",
22 "I'm planning a vacation.",
23 "She sings beautifully.",
24 "The cat is sleeping.",
25 "I want to learn French.",
26 "I enjoy going to the beach.",
27 "Where can I find a taxi?",
28 "I'm sorry for the inconvenience.",
29 "I'm studying for my exams.",
30 "I like to cook dinner at home.",
31 "Do you have any recommendations for restaurants?",
32 ]
33 espanhol = [
34 "Hola, cómo estás?",
35 "Me encanta leer libros.",
36 "El clima está agradable hoy.",
37 "Dónde está el restaurante más cercano?",
38 "Qué hora es?",
39 "Voy al parque todos los días.",
40 "Puedes ayudarme con esto?",
41 "Me gustaría ir de vacaciones.",
42 "Este es mi libro favorito.",
43 "Me gusta bailar salsa.",
44 "Hablas español?",
45 "Cuál es tu comida favorita?",
46 "Estoy aprendiendo a tocar el piano.",
47 "Que tengas un buen día!",
48 "Necesito comprar algunas frutas.",
49 "Vamos a dar un paseo.",
50 "Cómo estuvo tu fin de semana?",
51 "Estoy emocionado por el concierto.",
52 "Me pasas la sal, por favor?",
53 "Tengo una reunión a las 2 PM.",
54 "Estoy planeando unas vacaciones.",
55 "Ella canta hermosamente.",
56 "El perro está jugando.",
57 "Quiero aprender italiano.",
58 "Disfruto ir a la playa.",
59 "Dónde puedo encontrar un taxi?",
60 "Lamento las molestias.",
61 "Estoy estudiando para mis exámenes.",
62 "Me gusta cocinar la cena en casa.",
63 "Tienes alguna recomendación de restaurantes?",
64 ]
65 portugues = [
66 "Estou indo para o trabalho agora.",

```

```

67     "Adoro passar tempo com minha família.",
68     "Preciso comprar leite e pão.",
69     "Vamos ao cinema no sábado.",
70     "Gosto de praticar esportes ao ar livre.",
71     "O trânsito está terrível hoje.",
72     "A comida estava deliciosa!",
73     "Você já visitou o Rio de Janeiro?",
74     "Tenho uma reunião importante amanhã.",
75     "A festa começa às 20h.",
76     "Estou cansado depois de um longo dia de trabalho.",
77     "Vamos fazer um churrasco no final de semana.",
78     "O livro que estou lendo é muito interessante.",
79     "Estou aprendendo a cozinhar pratos novos.",
80     "Preciso fazer exercícios físicos regularmente.",
81     "Vou viajar para o exterior nas férias.",
82     "Você gosta de dançar?",
83     "Hoje é meu aniversário!",
84     "Gosto de ouvir música clássica.",
85     "Estou estudando para o vestibular.",
86     "Meu time de futebol favorito ganhou o jogo.",
87     "Quero aprender a tocar violão.",
88     "Vamos fazer uma viagem de carro.",
89     "O parque fica cheio aos finais de semana.",
90     "O filme que assisti ontem foi ótimo.",
91     "Preciso resolver esse problema o mais rápido possível.",
92     "Adoro explorar novos lugares.",
93     "Vou visitar meus avós no domingo.",
94     "Estou ansioso para as férias de verão.",
95     "Gosto de fazer caminhadas na natureza.",
96     "O restaurante tem uma vista incrível.",
97     "Vamos sair para jantar no sábado."
98 ]
99
100 import random
101 import re
102 import pandas as pd
103
104 pre_padroes = []
105 for frase in ingles:
106     pre_padroes.append( [frase, 'inglês'])
107 for frase in espanhol:
108     pre_padroes.append( [frase, 'espanhol'])
109
110 for frase in portuges:
111     pre_padroes.append( [frase, 'português'])
112 random.shuffle(pre_padroes)
113
114 dados = pd.DataFrame(pre_padroes)
115
116 def tamanhoMedioFrases(texto):
117     palavras = re.split("\s", texto)

```

```

118     tamanhos = [len(s) for s in palavras if len(s)>0]
119     #print(palavras)
120     #print(tamanhos)
121     soma = 0
122     for t in tamanhos:
123         soma=soma+t
124     return soma / len(tamanhos)
125
126 def verificaCaracteresEspeciais(frase):
127     # Criação de dicionário com regras de cada idioma
128     dictRegex = {
129         'nTio': r'[\~n ]',
130         'exclamacaoInvertido': r'[\textexclamdown]',
131         'interrogacaoInvertido': r'[\textquestiondown]',
132         'cedilhas': r'[ç]',
133         'vogalAcentoAgudo': r'[áéíóú]',
134         'vogalAcentoCrase': r'[à]',
135         'vogalAcentoCircunflexo': r'[âêô]',
136         'digrafos': r'ch|lh|nh|gu|qu|rr|ss|sc|sç|xc|ll|nn|cc',
137         'digrafosVogal': 'ae|ai|ao|au|ea|ei|eo|eu|ia|ie|io|iu|oa|oe|oi|
            ou|ua|ue|ui|uo',
138         'repeticaoVogal': r'aa|ee|ii|oo',
139         'ditongos': r'âe|ãe|õe',
140         'apostrofe': r"[']",
141         'formatoHoras': r'PM|AM',
142         'consoantesReservadas': r'[kwxy]'
143     }
144
145 def qtdOcorrencias(chave): return len(re.findall(dictRegex[chave],
    frase, re.IGNORECASE))
146 dictQtdCaracteresEspeciais = {
147     'qtdNtio': qtdOcorrencias('nTio'),
148     'qtdExclamacaoInvertido': qtdOcorrencias('exclamacaoInvertido'),
149     'qtdInterrogacaoInvertido': qtdOcorrencias('interrogacaoInvertido'
    ),
150     'qtdChedilhas': qtdOcorrencias('cedilhas'),
151     'qtdVogalAcentoCrase': qtdOcorrencias('vogalAcentoCrase'),
152     'qtdVogalAcentoAgudo': qtdOcorrencias('vogalAcentoAgudo'),
153     'qtdVogalAcentoCircunflexo': qtdOcorrencias('
    vogalAcentoCircunflexo'),
154     'qtdDigrafos': qtdOcorrencias('digrafos'),
155     'qtdDigrafosVogal': qtdOcorrencias('digrafosVogal'),
156     'qtdDitongos': qtdOcorrencias('ditongos'),
157     'qtdApostrofe': qtdOcorrencias('apostrofe'),
158     'qtdRepeticaoVogal': qtdOcorrencias('repeticaoVogal'),
159     'qtdHoras': qtdOcorrencias('formatoHoras'),
160     'qtdConsoantesReservadas': qtdOcorrencias('consoantesReservadas')
161 }
162 print(frase, '-', dictQtdCaracteresEspeciais.values())
163 return dictQtdCaracteresEspeciais.values()
164

```

```

165 def extraCaracteristicas(frase):
166     texto = frase[0]
167     pattern_regex = re.compile('[^\w+]', re.UNICODE)
168     texto = re.sub(pattern_regex, ' ', texto)
169     caracteristica1=tamanhoMedioFrases(texto)
170     caracteristica2=verificaCaracteresEspeciais(frase[0])
171     padrao = []
172     padrao.append(caracteristica1)
173     padrao.extend(caracteristica2)
174     padrao.append(frase[1])
175     return padrao
176
177 def geraPadroes(frases):
178     padroes = []
179     for frase in frases:
180         padrao = extraCaracteristicas(frase)
181         padroes.append(padrao)
182     return padroes
183
184 padroes = geraPadroes(pre_padroes)
185 dados = pd.DataFrame(padroes)
186
187 from sklearn.model_selection import train_test_split
188 import numpy as np
189
190 vet = np.array(padroes)
191 classes = vet[:, -1] # classes = [p[-1] for p in padroes]
192 padroes_sem_classe = vet[:, 0:-1]
193
194 X_train, X_test, y_train, y_test = train_test_split(
195     padroes_sem_classe, classes, test_size=0.25, stratify=classes)
196
197 from sklearn import svm
198 from sklearn.metrics import confusion_matrix
199 from sklearn.metrics import classification_report
200 treinador = svm.SVC() #algoritmo escolhido
201 modelo = treinador.fit(X_train, y_train)
202
203 acuracia = modelo.score(X_train, y_train)
204 print("Acurácia nos dados de treinamento: {:.2f}%".format(acuracia *
205     100))
206
207 # melhor avaliar com a matriz de confusão
208 y_pred = modelo.predict(X_train)
209 cm = confusion_matrix(y_train, y_pred)
210
211 print(classification_report(y_train, y_pred))
212
213 # score com os dados de teste
214 acuracia = modelo.score(X_test, y_test)
215 print("Acurácia nos dados de teste: {:.2f}%".format(acuracia * 100))

```

```

214
215 print('métricas mais confiáveis')
216 y_pred2 = modelo.predict(X_test)
217 cm = confusion_matrix(y_test, y_pred2)
218 print(cm)
219 print(classification_report(y_test, y_pred2))

```

Acurácia nos dados de treinamento: 75.36%

[[12 4 6]

[0 21 2]

[5 0 19]]

precision recall f1-score support

espanhol 0.71 0.55 0.62 22

inglês 0.84 0.91 0.87 23

português 0.70 0.79 0.75 24

accuracy 0.75 69

macro avg 0.75 0.75 0.75 69

weighted avg 0.75 0.75 0.75 69

Acurácia nos dados de teste: 82.61%

métricas mais confiáveis

[[5 1 2]

[1 6 0]

[0 0 8]]

precision recall f1-score support

espanhol 0.83 0.62 0.71 8

inglês 0.86 0.86 0.86 7

português 0.80 1.00 0.89 8

accuracy 0.83 23

macro avg 0.83 0.83 0.82 23

weighted avg 0.83 0.83 0.82 23

2 Melhore uma base de dados ruim

```

1 import requests, zipfile, io
2 from io import BytesIO
3 import numpy as np
4 import pandas as pd
5 from sklearn import svm
6 from sklearn.metrics import confusion_matrix
7 from sklearn.metrics import classification_report
8 from sklearn.preprocessing import LabelEncoder
9 # link dos dados: https://www.archive.ics.uci.edu/dataset/14/breast+
  cancer
10
11 data = pd.read_csv('/content/breast-cancer.data', header=None,
  lineterminator='\n')
12 data.columns = ['class', 'age', 'menopause', 'tumor-size', 'inv-nodes',
  'node-caps', 'deg-malig', 'breast', 'breast-quad', 'irradiat']

```

Por meio da análise dos dados presentes em cada coluna onde é possível tirar as seguintes conclusões:

- Os dados possuem natureza categórica
- Temos um desbalanceamento grande entre as classes
- Existe pouca prevalência de pessoas mais novas e de pessoas muito mais velhas
- Maior parte dos casos estão na pré-menopausa
- Prevalência de inv-nodes discrepantes aos demais casos
- Ocorrências de câncer nas mamas esquerda e direita são aproximadas
- É possível perceber a presença do carácter ? em alguns campos

```

1 # transformando os valores presentes nas colunas em inteiros
  fatorizados
2 labelEncoder = LabelEncoder()
3 for col in data.columns: data[col] = labelEncoder.fit_transform(data[
  col])
4
5 new_data = pd.read_csv('/content/breast-cancer.data', header=None,
  lineterminator='\n')
6 new_data.columns = data.columns
7
8 # remoção das linhas que possuem o carácter ?
9 new_data_clean = new_data[~new_data.apply(lambda row: row.astype(str)
  .str.contains('\?').any(), axis=1)]
10
11 new_data = new_data_clean.copy()
12 labelEncoder = LabelEncoder()
13 for col in new_data.columns: new_data[col] = labelEncoder.
  fit_transform(new_data[col])
14
15 # ao remover as colunas 'age', 'tumor-size e 'breast-quad' os valores
  do classificador melhoraram
16 # a conclusão acerca disso é que o atributo idade possui uma grande
  variedade de intervalos, o que pode dificultar a separação dos
  grupos
17 # o mesmo acontece para o tamanho do tumor. Em relação ao local (
  quadrante) onde o nódulo não foi significativo para a classificaçã
  o
18 # pois a informação sobre qual mama o nódulo já está presente na
  coluna 'breast', essa nova coluna só inclui o local mais preciso.
19
20 new_data = new_data.drop(columns=['age', 'tumor-size', 'breast-quad'])
21
22 X_orig = data.iloc[:, 1:].values
23 Y_orig = data.iloc[:, 0].values
24
25 X = new_data.iloc[:, 1:].values
26 Y = new_data.iloc[:, 0].values
27

```

```

28 from sklearn.model_selection import train_test_split
29 import numpy as np
30
31 # com os dados originais
32 X_orig_train, X_orig_test, y_orig_train, y_orig_test =
    train_test_split(X_orig,
33 Y_orig, test_size=0.25, stratify=Y_orig, random_state=10)
34
35 # com os dados tratados
36 X_train, X_test, y_train, y_test = train_test_split(X, Y, test_size
    =0.25,
37 stratify=Y, random_state=10)
38
39 # VERIFICAÇÃO DO MODELO **SEM** TRATAMENTO DOS DADOS
40
41 treinador = svm.SVC() #algoritmo escolhido
42 modelo_orig = treinador.fit(X_orig_train, y_orig_train)
43
44 # predição com os mesmos dados usados para treinar
45 print('Matriz de confusão - com os dados ORIGINAIS usados no
    TREINAMENTO')
46 print('Acurácia:', "{:.2f}%".format(modelo_orig.score(X_orig_train,
    y_orig_train)*100))
47
48 y_orig_pred = modelo_orig.predict(X_orig_train)
49 cm_orig_train = confusion_matrix(y_orig_train, y_orig_pred)
50 print(cm_orig_train)
51 print(classification_report(y_orig_train, y_orig_pred))
52
53 # predição com os mesmos dados usados para testar
54 print('Matriz de confusão - com os dados ORIGINAIS usados para TESTES
    ')
55 print('Acurácia:', "{:.2f}%".format(modelo_orig.score(X_orig_test,
    y_orig_test)*100))
56 y2_orig_pred = modelo_orig.predict(X_orig_test)
57 cm_orig_test = confusion_matrix(y_orig_test, y2_orig_pred)
58 print(cm_orig_test)
59 print(classification_report(y_orig_test, y2_orig_pred))

```

Matriz de confusão - com os dados ORIGINAIS usados no TREINAMENTO

Acurácia: 76.17%

[[146 4]

[47 17]]

precision recall f1-score support

0 0.76 0.97 0.85 150

1 0.81 0.27 0.40 64

accuracy 0.76 214

macro avg 0.78 0.62 0.63 214

weighted avg 0.77 0.76 0.72 214

Matriz de confusão - com os dados ORIGINAIS usados para TESTES

Acurácia: 76.39%

```

[[47 4]
 [13 8]]
precision recall f1-score support
0 0.78 0.92 0.85 51
1 0.67 0.38 0.48 21
accuracy 0.76 72
macro avg 0.72 0.65 0.67 72
weighted avg 0.75 0.76 0.74 72

```

```

1 # VERIFICAÇÃO DO MODELO **COM** TRATAMENTO DOS DADOS
2
3 treinador = svm.SVC() #algoritmo escolhido
4 modelo = treinador.fit(X_train, y_train)
5 # predição com os mesmos dados usados para treinar
6 print('Matriz de confusão - com os dados TRATADOS usados no
      TREINAMENTO')
7 print('Acurácia:', "{:.2f}%".format(modelo.score(X_train, y_train)
      *100))
8 y_pred = modelo.predict(X_train)
9 cm_train = confusion_matrix(y_train, y_pred)
10 print(cm_train)
11 print(classification_report(y_train, y_pred))
12
13 # predição com os mesmos dados usados para testar
14 print('Matriz de confusão - com os dados ORIGINAIS usados para TESTES
      ')
15 print('Acurácia:', "{:.2f}%".format(modelo.score(X_test, y_test)*100)
      )
16 y2_pred = modelo.predict(X_test)
17 cm_test = confusion_matrix(y_test, y2_pred)
18 print(cm_test)
19 print(classification_report(y_test, y2_pred))

```

Matriz de confusão - com os dados TRATADOS usados no TREINAMENTO

Acurácia: 78.74%

```
[[142 4]
```

```
[ 40 21]]
```

```
precision recall f1-score support
```

```
0 0.78 0.97 0.87 146
```

```
1 0.84 0.34 0.49 61
```

```
accuracy 0.79 207
```

```
macro avg 0.81 0.66 0.68 207
```

```
weighted avg 0.80 0.79 0.75 207
```

Matriz de confusão - com os dados ORIGINAIS usados para TESTES

Acurácia: 81.43%

```
[[48 2]
```

```
[11 9]]
```

```
precision recall f1-score support
```

```
0 0.81 0.96 0.88 50
```

```
1 0.82 0.45 0.58 20
```

```
accuracy 0.81 70
```

macro avg 0.82 0.70 0.73 70
weighted avg 0.81 0.81 0.79 70

APÊNDICE 7 – APRENDIZADO DE MÁQUINA

A - ENUNCIADO

Para cada uma das tarefas abaixo (Classificação, Regressão etc.) e cada base de dados (Veículo, Diabetes etc.), fazer os experimentos com todas as técnicas solicitadas (KNN, RNA etc.) e preencher os quadros com as estatísticas solicitadas, bem como os resultados pedidos em cada experimento.

B - RESOLUÇÃO

Classificação - Veículos

Tabela 7.1: Veículos - Resultados por técnica com parâmetros, acurácia e matrizes de confusão.

Técnica	Parâmetro	Acurácia	Matriz de Confusão
SVM – CV	C=100, Sigma=0.01	0.8443	Reference Prediction bus opel saab van bus 40 0 1 0 opel 0 31 10 0 saab 0 10 31 0 van 3 1 1 39
RNA – CV	size=31, decay=0.7	0.8204	Reference Prediction bus opel saab van bus 41 0 0 0 opel 0 27 13 0 saab 0 13 30 0 van 2 0 0 39
SVM – Hold-out	C=1, Sigma=0.07	0.7485	Reference Prediction bus opel saab van bus 40 0 0 1 opel 0 20 13 0 saab 0 20 26 0 van 3 0 0 39
RF – Hold-out	mtry=10	0.7425	Reference Prediction bus opel saab van bus 40 0 1 0 opel 0 19 13 0 saab 0 20 26 0 van 3 3 3 39

Continua na próxima página

Continuação da Tabela 7.1

Técnica	Parâmetro	Acurácia	Matriz de Confusão
RF – CV	mtry=5	0.7305	Reference Prediction bus opel saab van bus 40 0 1 0 opel 0 18 14 0 saab 0 21 25 0 van 3 3 0 39
RNA – Hold-out	size=5, decay=0.1	0.6707	Reference Prediction bus opel saab van bus 35 1 0 0 opel 0 0 0 0 saab 2 9 43 5 van 6 2 0 34
KNN	k=1	0.6353	Reference Prediction bus opel saab van bus 39 1 3 2 opel 1 18 19 2 saab 7 23 15 0 van 0 4 0 36

Classificação - Diabetes

Tabela 7.2: Diabetes - Resultados por técnica com parâmetros, acurácia e matriz de confusão.

Técnica	Parâmetro	Acurácia	Matriz de Confusão
SVM – Hold-out	C=0.5, Sigma=0.14	0.7647	Reference Prediction neg pos neg 92 28 pos 8 25
RF – Hold-out	mtry=2	0.7647	Reference Prediction neg pos neg 86 22 pos 14 31
RF – CV	mtry=9	0.7647	Reference Prediction neg pos neg 83 19 pos 17 34
SVM – CV	C=2, Sigma=0.015	0.7582	Reference Prediction neg pos neg 91 28 pos 9 25

Continua na próxima página

Continuação da Tabela 7.2

Técnica	Parâmetro	Acurácia	Matriz de Confusão		
KNN	k=7	0.7273	Reference		
			Prediction	neg	pos
			neg	85	31
			pos	11	27
RNA – CV	size=21, decay=0.4	0.7059	Reference		
			Prediction	neg	pos
			neg	81	26
			pos	19	27
RNA – Hold-out	size=5, decay=0.1	0.6732	Reference		
			Prediction	neg	pos
			neg	78	28
			pos	22	25

Regressão - Admissão

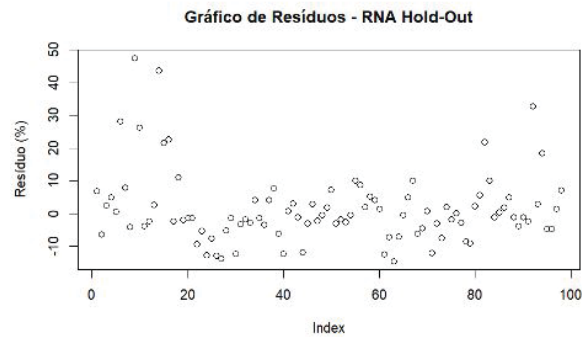
Tabela 7.3: Admissão - Resultados por técnica com parâmetros e métricas (R^2 , S_{yx} , coeficiente de Pearson, RMSE e MAE).

Técnica	Parâmetro	R^2	S_{yx}	Pearson	RMSE	MAE
RNA – Hold-out	size=3, decay=1e-4	0.8234	0.0606	0.9083	0.0583	0.0434
RF – Hold-out	mtry=2	0.8231	0.0606	0.9089	0.0584	0.0431
RF – CV	mtry=2	0.8196	0.0612	0.9072	0.0590	0.0436
SVM – CV	C=50, Sigma=0.01	0.8160	0.0618	0.9063	0.0595	0.0430
SVM – Hold-out	C=1, Sigma=0.18	0.8014	0.0642	0.9007	0.0619	0.0438
KNN	k=9	0.7384	0.0737	0.8600	0.0710	0.0547
RNA – CV	size=41, decay=0.1	0.7344	0.0743	0.8703	0.0716	0.0534

Tabela 7.4: Amostra com GRE, TOEFL, rating universitário e métricas preditivas.

GRE.Score	TOEFL.Score	University.Rating	SOP	LOR	CGPA	Research	predict.rna
316.99	110.72	4.66	1.19	2.20	9.70	0.71	0.89
303.68	99.41	2.67	4.11	2.01	8.08	0.41	0.59
295.93	97.60	4.72	3.48	2.84	9.80	0.80	0.82

Figura 7.1 – ADMISSÃO - GRÁFICO DE RESÍDUOS (MELHOR MODELO)



Regressão - Biomassa

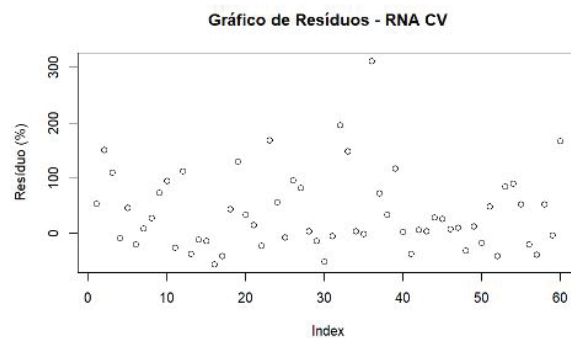
Tabela 7.5: Biomassa - Resultados por técnica com parâmetros e métricas (R^2 , $S_{y.x}$, coeficiente de Pearson, RMSE e MAE).

Técnica	Parâmetro	R^2	$S_{y.x}$	Pearson	RMSE	MAE
RNA – CV	size=11, decay=0.4	0.9877	141.04	0.9939	137.47	64.407
RF – CV	mtry=5	0.9846	163.73	0.9923	153.88	61.807
RF – Hold-out	mtry=2	0.9843	165.45	0.9921	155.50	63.321
SVM – CV	C=100, Sigma=0.01	0.9448	310.55	0.9881	291.87	94.531
KNN	k=3	0.8576	498.83	0.9578	468.83	124.52
RNA – Hold-out	size=5, decay=0.1	0.5433	861.62	0.8205	839.80	186.86
SVM – Hold-out	C=1, Sigma=0.60	0.3435	1071.3	0.6175	1006.89	210.16

Tabela 7.6: Amostra com variáveis *dap*, *h*, *Me* e predição da RNA.

dap	h	Me	predict.rna
132.31429	23.76053	0.4378101	16744.171
55.74442	48.14359	0.6876391	2724.243
45.74574	14.74192	0.3821269	14206.252

Figura 7.2 – BIOMASSA - GRÁFICO DE RESÍDUOS (MELHOR MODELO)



Agrupamento - Veículo

K-means clustering with 10 clusters of size 75, 67, 54, 131, 24, 137, 78, 109, 70, 101

Tabela 7.7: K-means (médias por cluster) - Parte 1

Cl	Comp	Circ	DCirc	RadRa	PrAxisRa	MaxLRa	ScatRa	Elong	PrAxisRect	MaxLRect
1	89.96	39.17	75.71	168.43	65.49	9.40	149.69	44.17	19.00	135.45
2	95.99	45.96	90.64	196.25	64.34	8.31	182.42	35.90	21.49	148.67
3	85.67	36.74	57.31	122.44	56.19	5.44	120.81	55.72	17.22	128.41
4	85.61	43.20	70.41	138.50	58.93	7.20	148.49	45.30	18.94	143.68
5	107.25	55.54	103.42	190.25	55.79	5.79	248.71	26.88	27.08	168.08
6	104.12	53.66	102.67	201.88	62.42	10.47	217.18	30.70	24.42	169.28
7	102.14	49.71	100.51	204.21	63.67	9.37	200.67	32.81	22.95	156.59
8	88.88	38.89	66.83	138.25	57.81	7.27	134.41	49.92	18.03	135.41
9	93.47	41.33	80.15	188.79	66.96	9.14	165.79	39.61	20.14	139.43
10	91.05	45.54	80.15	157.95	63.08	10.02	156.26	43.02	19.56	152.19

Tabela 7.8: K-means (médias por cluster) - Parte 2

Cl	ScVarMaxis	ScVarmaxis	RaGyr	SkewMaxis	Skewmaxis	Kurtmaxis	KurtMaxis	HolIRa
1	173.88	335.92	142.24	69.19	4.72	14.37	193.47	200.41
2	203.27	502.81	179.22	68.01	6.24	12.97	193.78	200.34
3	140.31	215.50	134.94	74.96	7.00	12.19	185.94	190.83
4	170.03	324.98	173.22	77.76	6.00	9.14	183.64	188.79
5	271.29	910.54	247.17	83.63	6.50	14.21	182.67	183.63
6	227.10	697.25	212.86	71.92	7.58	15.88	187.91	197.39
7	218.59	603.91	199.73	69.38	6.46	14.76	191.15	199.00
8	156.91	267.55	142.82	71.35	6.09	11.05	189.72	194.93
9	192.69	414.83	155.57	69.83	5.40	15.00	194.31	200.69
10	176.25	360.31	178.36	73.12	7.11	9.27	187.40	195.52

Associação - Musculação

Tabela 7.9: Regras de associação — top 10 linhas.

Lhs	Rhs	Support	Confidence	Coverage	Lift	Count
{V2=Gêmeos}	{V3=Bicicleta}	0.026	1.0	0.026	3.2	1
{V2=Gêmeos}	{V4=Esteira}	0.026	1.0	0.026	3.2	1
{V2=Gêmeos}	{V1=Extensor}	0.026	1.0	0.026	1.0	1
{V3=Gêmeos}	{V2=AgachamentoSmith}	0.051	1.0	0.051	3.9	2
{V3=Gêmeos}	{V4=Esteira}	0.051	1.0	0.051	3.2	2
{V4=Crucifixo}	{V3=Afundo}	0.051	1.0	0.051	6.5	2
{V4=Crucifixo}	{V2=Gêmeos}	0.051	1.0	0.051	5.5	2
{V4=Crucifixo}	{V1=LegPress}	0.051	1.0	0.051	2.6	2
{V2=Flexor}	{V3=Bicicleta}	0.051	1.0	0.051	3.2	2
{V2=Flexor}	{V4=Esteira}	0.051	1.0	0.051	3.2	2

Código

```

1  install.packages("e1071")
2  install.packages("caret")
3  install.packages("Metrics")
4
5  library("caret")
6  library("Metrics")
7
8  % CLASSIFICACAO - VEICULOS
9
10 setwd("/Users/blb/Documents/ESPECIALIZAÇÃO EM IA/APRENDIZADO DE
    MAQUINA/CLASSIFICACAO/")
11 dados <- read.csv("6 - Veiculos - Dados.csv")
12 dados[[a]] <- NULL
13
14 % KNN
15
16 set.seed(202410)
17 ran <- sample(1:nrow(dados), 0.8 * nrow(dados))
18 treino <- dados[ran,]
19 teste <- dados[-ran,]
20 set.seed(202410)
21 tuneGrid <- expand.grid(k = c(1,3,5,7,9))
22 knn <- train(tipo ~ ., data = treino, method = "knn", tuneGrid=
    tuneGrid)
23 predict.knn <- predict(knn, teste)
24 confusionMatrix(predict.knn, as.factor(teste[["tipo"]]))
25
26 % RNA HOLD OUT
27
28 set.seed(202410)
29 indices <- createDataPartition(dados[["tipo"]], p=0.80, list=FALSE)
30 treino <- dados[indices,]
31 teste <- dados[-indices,]
32 set.seed(202410)
33 rna <- train(tipo ~ ., data=treino, method="nnet", trace=FALSE)
34 predicoes.rna <- predict(rna, teste)
35 confusionMatrix(predicoes.rna, as.factor(teste[["tipo"]]))
36
37 % RNA - Cross-Validation
38
39 ctrl <- trainControl(method = "cv", number = 10)
40 grid <- expand.grid(size = seq(from = 1, to = 45, by = 10), decay =
    seq(from = 0.1, to = 0.9,
41 by = 0.3))
42 set.seed(202410)
43 rna <- train(
44   form = tipo~.,
45   data = treino,
46   method = "nnet",
47   tuneGrid = grid,
48   trControl = ctrl,

```

```

49     maxit = 2000,trace=FALSE)
50 predicoes.rna <- predict(rna, teste)
51 confusionMatrix(predicoes.rna, as.factor(teste[["tipo"]]))
52
53 % SVM - Hold-out
54
55 set.seed(202410)
56 indices <- createDataPartition(dados[["tipo"]], p=0.80, list=FALSE)
57 treino <- dados[indices,]
58 teste <- dados[-indices,]
59 set.seed(202410)
60 svm <- train(tipo~., data=treino, method="svmRadial")
61 predicoes.svm <- predict(svm, teste)
62 confusionMatrix(predicoes.svm, as.factor(teste[["tipo"]]))
63
64 % SVM - Cross-Validation
65
66 set.seed(202410)
67 indices <- createDataPartition(dados[["tipo"]], p=0.80, list=FALSE)
68 treino <- dados[indices,]
69 teste <- dados[-indices,]
70 ctrl <- trainControl(method = "cv", number = 10)
71 tuneGrid = expand.grid(C=c(1, 2, 10, 50, 100), sigma=c(.01, .015,
72     0.2))
73 set.seed(202410)
74 svm <- train(tipo~., data=treino, method="svmRadial", trControl=ctrl,
75     tuneGrid=tuneGrid)
76
77 predicoes.svm <- predict(svm, teste)
78 confusionMatrix(predicoes.svm, as.factor(teste[["tipo"]]))
79
80 % RANDOM FOREST - Hold-out
81
82 dados_aux <- dados
83 dados_aux[["tipo"]] <- NULL
84 newdata <- lapply(dados_aux, function(coluna) {
85     return(runif(3, min = min(coluna), max = max(coluna)))
86 })
87 predicoes.svm <- predict(svm, data.frame(newdata))
88 resultado <- cbind(data.frame(newdata), predicoes.svm)
89 View(resultado)
90
91 set.seed(202410)
92 indices <- createDataPartition(dados[["tipo"]], p=0.80, list=FALSE)
93 treino <- dados[indices,]
94 teste <- dados[-indices,]
95 set.seed(202410)
96 rf <- train(tipo~., data=treino, method="rf")
97
98 predicoes.rf <- predict(rf, teste)
99 confusionMatrix(predicoes.rf, as.factor(teste[["tipo"]]))

```

```

98 % RANDOM FOREST - Cross-Validation
99
100 set.seed(202410)
101 indices <- createDataPartition(dados[["tipo"]], p=0.80, list=FALSE)
102 treino <- dados[indices,]
103 teste <- dados[-indices,]
104 ctrl <- trainControl(method = "cv", number = 10)
105 tuneGrid = expand.grid(mtry=c(2, 5, 7, 9))
106 set.seed(202410)
107 rf <- train(tipo~., data=treino, method="rf", trControl=ctrl,
108             tuneGrid=tuneGrid)
109 predicoes.rf <- predict(rf, teste)
110 confusionMatrix(predicoes.rf, as.factor(teste[["tipo"]]))

```

```

1 % CLASSIFICACAO - DIABETES
2
3 setwd("/Users/blb/Documents/ESPECIALIZAÇÃO EM IA/APRENDIZADO DE
  MAQUINA/CLASSIFICACAO/")
4 dados <- read.csv("10 - Diabetes - Dados.csv")
5 dados[["num"]] <- NULL
6
7 % KNN
8
9 set.seed(202410)
10 ran <- sample(1:nrow(dados), 0.8 * nrow(dados))
11 treino <- dados[ran,]
12 teste <- dados[-ran,]
13 set.seed(202410)
14 tuneGrid <- expand.grid(k = c(1,3,5,7,9))
15 knn <- train(diabetes ~ ., data = treino, method = "knn", tuneGrid=
16             tuneGrid)
17 predict.knn <- predict(knn, teste)
18 confusionMatrix(predict.knn, as.factor(teste[["diabetes"]]))
19
20 % RNA - Hold-out
21
22 set.seed(202410)
23 indices <- createDataPartition(dados[["diabetes"]], p=0.80, list=
24                                 FALSE)
25 treino <- dados[indices,]
26 teste <- dados[-indices,]
27 set.seed(202410)
28 rna <- train(diabetes ~ ., data=treino, method="nnet", trace=FALSE)
29 predicoes.rna <- predict(rna, teste)
30 confusionMatrix(predicoes.rna, as.factor(teste[["diabetes"]]))
31
32 % RNA - Cross-Validation
33
34 ctrl <- trainControl(method = "cv", number = 10)
35 grid <- expand.grid(size = seq(from = 1, to = 45, by = 10), decay =
36                     seq(from = 0.1, to = 0.9,

```

```

34 by = 0.3))
35 set.seed(202410)
36 rna <- train(
37   form = diabetes~.,
38   data = treino,
39   method = "nnet",
40   tuneGrid = grid,
41   trControl = ctrl,
42   maxit = 2000, trace=FALSE)
43 predicoes.rna <- predict(rna, teste)
44 confusionMatrix(predicoes.rna, as.factor(teste[["diabetes"]]))
45
46 % SVM - Hold-out
47
48 set.seed(202410)
49 indices <- createDataPartition(dados[["diabetes"]], p=0.80, list=
50   FALSE)
51 treino <- dados[indices,]
52 teste <- dados[-indices,]
53 set.seed(202410)
54 svm <- train(diabetes~., data=treino, method="svmRadial")
55 predicoes.svm <- predict(svm, teste)
56 confusionMatrix(predicoes.svm, as.factor(teste[["diabetes"]]))
57
58 dados_aux <- dados
59 dados_aux[["diabetes"]] <- NULL
60 newdata <- lapply(dados_aux, function(coluna) {
61   return(runif(3, min = min(coluna), max = max(coluna)))
62 })
63 predicoes.svm <- predict(svm, data.frame(newdata))
64 resultado <- cbind(data.frame(newdata), predicoes.svm)
65
66 % SVM - Cross-Validation
67
68 set.seed(202410)
69 indices <- createDataPartition(dados[["diabetes"]], p=0.80, list=
70   FALSE)
71 treino <- dados[indices,]
72 teste <- dados[-indices,]
73 ctrl <- trainControl(method = "cv", number = 10)
74 tuneGrid = expand.grid(C=c(1, 2, 10, 50, 100), sigma=c(.01, .015,
75   0.2))
76 set.seed(202410)
77 svm <- train(diabetes~., data=treino, method="svmRadial", trControl=
78   ctrl, tuneGrid=tuneGrid)
79 predicoes.svm <- predict(svm, teste)
80 confusionMatrix(predicoes.svm, as.factor(teste[["diabetes"]]))
81
82 % RANDOM FOREST - Hold-out
83
84 set.seed(202410)

```

```

81 indices <- createDataPartition(dados[["diabetes"]], p=0.80, list=
    FALSE)
82 treino <- dados[indices,]
83 teste <- dados[-indices,]
84 set.seed(202410)
85 rf <- train(diabetes~., data=treino, method="rf")
86 predicoes.rf <- predict(rf, teste)
87 confusionMatrix(predicoes.rf, as.factor(teste[["diabetes"]]))
88
89 % RANDOM FOREST - Cross-Validation
90
91 set.seed(202410)
92 indices <- createDataPartition(dados[["diabetes"]], p=0.80, list=
    FALSE)
93 treino <- dados[indices,]
94 teste <- dados[-indices,]
95 ctrl <- trainControl(method = "cv", number = 10)
96 tuneGrid = expand.grid(mtry=c(2, 5, 7, 9))
97 set.seed(202410)
98 rf <- train(diabetes~., data=treino, method="rf", trControl=ctrl,
    tuneGrid=tuneGrid)
99 predicoes.rf <- predict(rf, teste)
100 confusionMatrix(predicoes.rf, as.factor(teste[["diabetes"]]))

```

```

1 % REGRESSAO - ADMISSAO
2
3 setwd("/Users/blb/Documents/ESPECIALIZAÇÃO EM IA/APRENDIZADO DE MÁ
    QUINA/REGRESSÃO/")
4 dados <- read.csv("9 - Admissao - Dados.csv", header=T)
5 dados[["num"]] <- NULL
6
7 r2 <- function(predito, observado) {
8 return(1 - (sum((predito-observado)^2) / sum((observado-mean(
    observado))^2)))
9 }
10
11 syx <- function(predito, observado, num_variaveis) {
12 return(sqrt((sum((predito-observado)^2) / (length(predito) -
    num_variaveis))))
13 }
14
15 regression_metrics <- function(predito, observado){
16 r2_aux <- r2(predito, observado)
17 syx_aux <- syx(predito, observado, 7)
18 pearson_aux <- cor(predito, observado, method='pearson')
19 rsme_aux <- rmse(predito, observado)
20 mae_aux <- mae(predito, observado)
21 return(c(r2_aux, syx_aux, pearson_aux, rsme_aux, mae_aux))
22 }
23
24 set.seed(202410)

```

```

25 ind <- createDataPartition(dados[["ChanceOfAdmit"]], p=0.80, list =
    FALSE)
26 treino <- dados[ind,]
27 teste <- dados[-ind,]
28 tuneGrid <- expand.grid(k = c(1,3,5,7,9))
29 set.seed(202410)
30 knn <- train(ChanceOfAdmit ~ ., data = treino, method = "knn",
    tuneGrid=tuneGrid)
31 predict.knn <- predict(knn, teste)
32 regression_metrics(predict.knn, teste[["ChanceOfAdmit"]])
33
34 % RNA - Hold-out
35
36 set.seed(202410)
37 indices <- createDataPartition(dados[["ChanceOfAdmit"]], p=0.80, list
    =FALSE)
38 treino <- dados[indices,]
39 teste <- dados[-indices,]
40 set.seed(202410)
41 rna <- train(ChanceOfAdmit ~ ., data=treino, method="nnet", trace=
    FALSE)
42 predict.rna <- predict(rna, teste)
43 regression_metrics(predict.rna, teste[["ChanceOfAdmit"]])
44 residuos <- function(predito, observado){
45   return(((predito - observado)/observado) * 100)
46 }
47 plot(residuos(predict.rna, teste[["ChanceOfAdmit"]]), ylab='Resíduo
    (%)', main='Gráfico de Resíduo
48 s - RNA Hold-Out')
49
50 dados_aux <- dados
51 dados_aux[["ChanceOfAdmit"]] <- NULL
52 newdata <- lapply(dados_aux, function(coluna) {
53   return(runif(3, min = min(coluna), max = max(coluna)))
54 })
55
56 predict.rna <- predict(rna, data.frame(newdata))
57 resultado <- cbind(data.frame(newdata), predict.rna)
58
59 % RNA - Cross-Validation
60
61 ctrl <- trainControl(method = "cv", number = 10)
62 grid <- expand.grid(size = seq(from = 1, to = 45, by = 10), decay =
    seq(from = 0.1, to = 0.9,
63 by = 0.3))
64 set.seed(202410)
65 rna <- train(
66   form = ChanceOfAdmit~.,
67   data = treino,
68   method = "nnet",
69   tuneGrid = grid,

```

```

70   trControl = ctrl,
71   maxit = 2000, trace=FALSE)
72
73 predict.rna <- predict(rna, teste)
74 regression_metrics(predict.rna, teste[["ChanceOfAdmit"]])
75
76 % SVM - Hold-out
77
78 set.seed(202410)
79 indices <- createDataPartition(dados[["ChanceOfAdmit"]], p=0.80, list
   =FALSE)
80 treino <- dados[indices,]
81 teste <- dados[-indices,]
82 set.seed(202410)
83 svm <- train(ChanceOfAdmit~., data=treino, method="svmRadial")
84 predict.svm <- predict(svm, teste)
85 regression_metrics(predict.svm, teste[["ChanceOfAdmit"]])
86
87 % SVM - Cross-Validation
88
89 set.seed(202410)
90 indices <- createDataPartition(dados[["ChanceOfAdmit"]], p=0.80, list
   =FALSE)
91 treino <- dados[indices,]
92 teste <- dados[-indices,]
93 ctrl <- trainControl(method = "cv", number = 10)
94 tuneGrid = expand.grid(C=c(1, 2, 10, 50, 100), sigma=c(.01, .015,
   0.2))
95 set.seed(202410)
96 svm <- train(ChanceOfAdmit~., data=treino, method="svmRadial",
   trControl=ctrl, tuneGrid=tuneG
97 rid)
98 predict.svm <- predict(svm, teste)
99 regression_metrics(predict.svm, teste[["ChanceOfAdmit"]])
100
101 % RANDOM FOREST - Hold-out
102
103 set.seed(202410)
104 indices <- createDataPartition(dados[["ChanceOfAdmit"]], p=0.80, list
   =FALSE)
105 treino <- dados[indices,]
106 teste <- dados[-indices,]
107 set.seed(202410)
108 rf <- train(ChanceOfAdmit~., data=treino, method="rf")
109 predict.rf <- predict(rf, teste)
110 regression_metrics(predict.rf, teste[["ChanceOfAdmit"]])
111
112 % RANDOM FOREST - Cross-Validation
113
114 set.seed(202410)

```

```

115 indices <- createDataPartition(dados[["ChanceOfAdmit"]], p=0.80, list
    =FALSE)
116 treino <- dados[indices,]
117 teste <- dados[-indices,]
118 ctrl <- trainControl(method = "cv", number = 10)
119 tuneGrid = expand.grid(mtry=c(2, 5, 7, 9))
120 set.seed(202410)
121 rf <- train(ChanceOfAdmit~., data=treino, method="rf", trControl=ctrl
    , tuneGrid=tuneGrid)
122 predict.rf <- predict(rf, teste)
123 regression_metrics(predict.rf, teste[["ChanceOfAdmit"]])

```

```

1 % REGRESSAO - BIOMASSA
2
3 setwd("/Users/blb/Documents/ESPECIALIZAÇÃO EM IA/APRENDIZADO DE MÁ
    QUINA/REGRESSÃO/")
4 dados <- read.csv("5 - Biomassa - Dados.csv", header=T)
5 dados[["num"]] <- NULL
6
7 % KNN
8
9 r2 <- function(predito, observado) {
10 return(1 - (sum((predito-observado)^2) / sum((observado-mean(
    observado))^2)))
11 }
12 syx <- function(predito, observado, num_variaveis) {
13 return(sqrt((sum((predito-observado)^2) / (length(predito) -
    num_variaveis))))
14 }
15
16 regression_metrics <- function(predito, observado){
17 r2_aux <- r2(predito, observado)
18 syx_aux <- syx(predito, observado, 3)
19 pearson_aux <- cor(predito, observado, method='pearson')
20 rsme_aux <- rmse(predito, observado)
21 mae_aux <- mae(predito, observado)
22 return(c(r2_aux, syx_aux, pearson_aux, rsme_aux, mae_aux))
23 }
24
25 residuos <- function(predito, observado){
26 return(((predito - observado)/observado) * 100)
27 }
28
29 set.seed(202410)
30 ind <- createDataPartition(dados[["biomassa"]], p=0.80, list = FALSE)
31 treino <- dados[ind,]
32 teste <- dados[-ind,]
33 tuneGrid <- expand.grid(k = c(1,3,5,7,9))
34 set.seed(202410)
35 knn <- train(biomassa ~ ., data = treino, method = "knn", tuneGrid=
    tuneGrid)

```

```

36 predict.knn <- predict(knn, teste)
37 regression_metrics(predict.knn, teste[["biomassa"]])
38
39
40 % RNA - Hold-out
41
42 set.seed(202410)
43 indices <- createDataPartition(dados[["biomassa"]], p=0.80, list=
  FALSE)
44 treino <- dados[indices,]
45 teste <- dados[-indices,]
46 set.seed(202410)
47 rna <- train(biomassa ~ ., data=treino, method="nnet", trace=FALSE,
  linout=T)
48 predict.rna <- predict(rna, teste)
49 regression_metrics(predict.rna, teste[["biomassa"]])
50
51 % RNA - Cross-Validation
52
53 ctrl <- trainControl(method = "cv", number = 10)
54 grid <- expand.grid(size = seq(from = 1, to = 45, by = 10), decay =
  seq(from = 0.1, to = 0.9,
55 by = 0.3))
56 set.seed(202410)
57 rna <- train(
58 form = biomassa~.,
59 data = treino,
60 method = "nnet",
61 tuneGrid = grid,
62 trControl = ctrl,
63 maxit = 2000, trace=FALSE, linout=T)
64 predict.rna <- predict(rna, teste)
65 regression_metrics(predict.rna, teste[["biomassa"]])
66
67 dap = runif(3, min = min(dados[["dap"]]), max = max(dados[["dap"]]))
68 h = runif(3, min = min(dados[["h"]]), max = max(dados[["h"]]))
69 Me = runif(3, min = min(dados[["Me"]]), max = max(dados[["Me"]]))
70 newdata <- data.frame(dap=dap, h=h, Me=Me)
71 predict.rna <- predict(rna, newdata)
72 resultado <- cbind(newdata, predict.rna)
73
74 % SVM - Hold-out
75
76 set.seed(202410)
77 indices <- createDataPartition(dados[["biomassa"]], p=0.80, list=
  FALSE)
78 treino <- dados[indices,]
79 teste <- dados[-indices,]
80 set.seed(202410)
81 svm <- train(biomassa~., data=treino, method="svmRadial")
82 predict.svm <- predict(svm, teste)

```

```

83 regression_metrics(predict.svm, teste[["biomassa"]])
84
85 % SVM - Cross-Validation
86
87 set.seed(202410)
88 indices <- createDataPartition(dados[["biomassa"]], p=0.80, list=
  FALSE)
89 treino <- dados[indices,]
90 teste <- dados[-indices,]
91 ctrl <- trainControl(method = "cv", number = 10)
92 tuneGrid = expand.grid(C=c(1, 2, 10, 50, 100), sigma=c(.01, .015,
  0.2))
93 set.seed(202410)
94 svm <- train(biomassa~., data=treino, method="svmRadial", trControl=
  ctrl, tuneGrid=tuneGrid)
95 predict.svm <- predict(svm, teste)
96 regression_metrics(predict.svm, teste[["biomassa"]])
97
98 % RANDOM FOREST - Hold-out
99
100 set.seed(202410)
101 indices <- createDataPartition(dados[["biomassa"]], p=0.80, list=
  FALSE)
102 treino <- dados[indices,]
103 teste <- dados[-indices,]
104 set.seed(202410)
105 rf <- train(biomassa~., data=treino, method="rf")
106 predict.rf <- predict(rf, teste)
107 regression_metrics(predict.rf, teste[["biomassa"]])
108
109 % RANDOM FOREST - Cross-Validation
110
111 set.seed(202410)
112 indices <- createDataPartition(dados[["biomassa"]], p=0.80, list=
  FALSE)
113 treino <- dados[indices,]
114 teste <- dados[-indices,]
115 ctrl <- trainControl(method = "cv", number = 10)
116 tuneGrid = expand.grid(mtry=c(2, 5, 7, 9))
117 set.seed(202410)
118 rf <- train(biomassa~., data=treino, method="rf", trControl=ctrl,
  tuneGrid=tuneGrid, linout=
119 T)
120 predict.rf <- predict(rf, teste)
121 regression_metrics(predict.rf, teste[["biomassa"]])

1 % AGRUPAMENTO - VEICULOS
2
3 setwd("/Users/blb/Documents/ESPECIALIZAÇÃO EM IA/APRENDIZADO DE MÁ
  QUINA/AGRUPAMENTO/")
4 dados <- read.csv("6 - Veiculos - Dados.csv")

```

```
5 dados[["a"]] <- NULL
6 dados[["tipo"]] <- NULL
7
8 set.seed(202410)
9 dados_normalizados <- scale(dados)
10 km.res = kmeans(dados, 10)
```

```
1 % REGRAS DE ASSOCIACAO - MUSCULACAO
2
3 install.packages('arules', dep=T)
4 library(arules)
5 library(datasets)
6
7 setwd("/Users/blb/Documents/ESPECIALIZAÇÃO EM IA/APRENDIZADO DE MÁ
  QUINA/ASSOCIAÇÃO/")
8 dados <- read.csv("2 - Musculacao - Dados(in).csv", header=0, sep=';')
9
10 set.seed(202410)
11 rules <- apriori(dados, parameter = list(supp = 0.001, conf = 0.7,
  minlen=2))
12 summary(rules)
13 options(digits=2)
14 inspect(sort(rules[1:20], by="confidence"))
```

APÊNDICE 8 – DEEP LEARNING

A - ENUNCIADO

1 Classificação de Imagens (CNN)

Implementar o exemplo de classificação de objetos usando a base de dados CIFAR10 e a arquitetura CNN vista no curso.

2 Detector de SPAM (RNN)

Implementar o detector de spam visto em sala, usando a base de dados SMS Spam e arquitetura de RNN vista no curso.

3 Gerador de Dígitos Fake (GAN)

Implementar o gerador de dígitos fake usando a base de dados MNIST e arquitetura GAN vista no curso.

4 Tradutor de Textos (Transformer)

Implementar o tradutor de texto do português para o inglês, usando a base de dados e a arquitetura Transformer vista no curso.

B - RESOLUÇÃO

1 Classificação de Imagens

```
1 import tensorflow as tf
2 import numpy as np
3 import matplotlib.pyplot as plt
4 from tensorflow.keras.layers import Input, Conv2D, Dense, Flatten,
    Dropout
5 from tensorflow.keras.models import Model
6 from mlxtend.plotting import plot_confusion_matrix
7 from sklearn.metrics import confusion_matrix
8
9 cifar10 = tf.keras.datasets.cifar10
10 (x_train, y_train), (x_test, y_test) = cifar10.load_data()
11
12 x_train, x_test = x_train/255.0, x_test/255.0
13 y_train, y_test = y_train.flatten(), y_test.flatten()
14
15 print("x_train.shape", x_train.shape)
16 print("y_train.shape", y_train.shape)
17 print("x_test.shape", x_test.shape)
18 print("y_test.shape", y_test.shape)
```

```
x_train.shape (50000, 32, 32, 3)
y_train.shape (50000,)
x_test.shape (10000, 32, 32, 3)
y_test.shape (10000,)
```

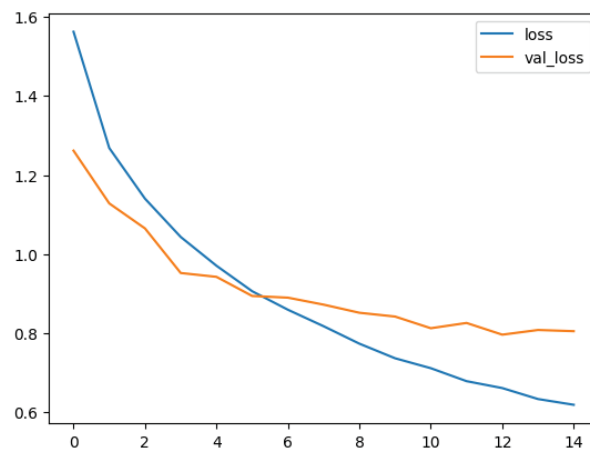
```
1 K = len(set(y_train))
2
3 i = Input(shape=x_train[0].shape)
4
5 x = Conv2D(32, (3,3), strides=2, activation="relu")(i)
```

```

6 x = Conv2D(64, (3,3), strides=2, activation="relu")(x)
7 x = Conv2D(128, (3,3), strides=2, activation="relu")(x)
8
9 x = Flatten()(x)
10 x = Dropout(0.5)(x)
11 x = Dense(1024, activation="relu")(x)
12 x = Dropout(0.2)(x)
13 x = Dense(K, activation="softmax")(x)
14 model = Model(i, x)
15
16 model.summary()
17
18 model.compile(optimizer="adam", loss="sparse_categorical_crossentropy",
19               metrics=["accuracy"])
20 r = model.fit(x_train, y_train, validation_data=(x_test, y_test),
21             epochs=15)
22 plt.plot(r.history["loss"], label="loss")
23 plt.plot(r.history["val_loss"], label="val_loss")
24 plt.legend()
25 plt.show()

```

Figura 8.1 – CLASSIFICAÇÃO DE IMAGENS COM CNN: GRÁFICO DE PERDA

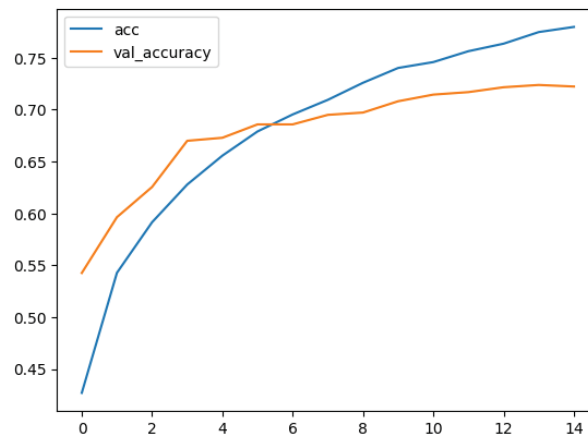


```

1 plt.plot(r.history["accuracy"], label="acc")
2 plt.plot(r.history["val_accuracy"], label="val_accuracy")
3 plt.legend()
4 plt.show()

```

Figura 8.2 – CLASSIFICAÇÃO DE IMAGENS COM CNN: GRÁFICO DE ACURÁCIA

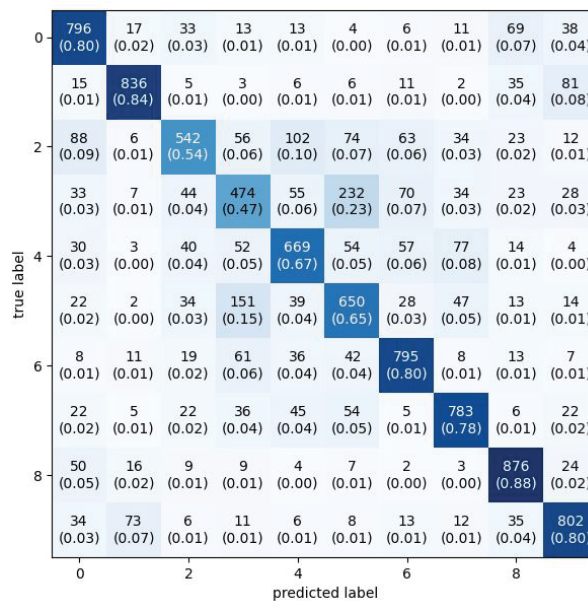


```

1 y_pred = model.predict(x_test).argmax(axis=1)
2 cm = confusion_matrix(y_test, y_pred)
3 plot_confusion_matrix(conf_mat=cm, figsize=(7,7), show_normed=True)

```

Figura 8.3 – CLASSIFICAÇÃO DE IMAGENS COM CNN: MATRIZ DE CONFUSÃO

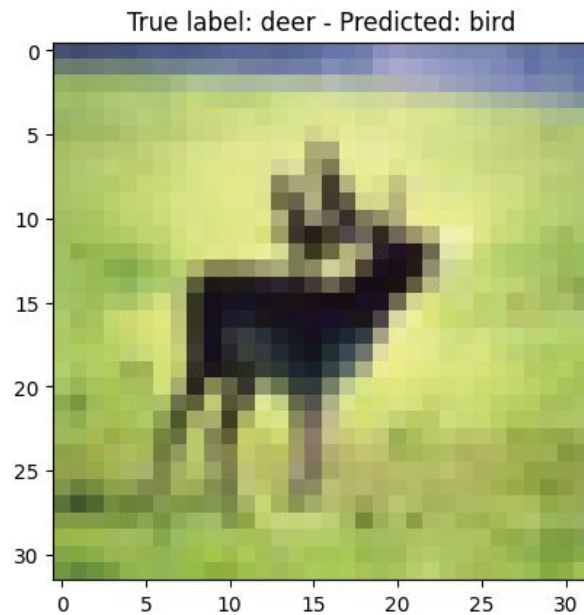


```

1
2 abels = ["airplane", "automobile", "bird", "cat", "deer", "dog", "frog", "horse", "ship", "truck"]
3
4 misclassified = np.where(y\_pred != y\_test)[0]
5
6 i = np.random.choice(misclassified)
7
8 plt.imshow(x\_test[i], cmap="gray")
9 plt.title("True label: \%s - Predicted: \%s" \% (labels[y\_test[i]], labels[y\_pred[i]]))

```

Figura 8.4 – CLASSIFICAÇÃO DE IMAGENS COM CNN: RESULTADO DA PREDIÇÃO



2 Detector de SPAM

```

1 import tensorflow as tf
2 import numpy as np
3 import matplotlib.pyplot as plt
4 import pandas as pd
5 from sklearn.model_selection import train_test_split
6 from tensorflow.keras.layers import Input, Embedding, LSTM, Dense
7 from tensorflow.keras.layers import GlobalMaxPooling1D
8 from tensorflow.keras.models import Model
9 from tensorflow.keras.preprocessing.sequence import pad_sequences
10 from tensorflow.keras.preprocessing.text import Tokenizer
11 from google.colab import files
12 import io
13
14 uploaded = files.upload()
15
16 df = pd.read_csv(io.BytesIO(uploaded['spam.csv']), encoding="ISO
    -8859-1")
17 df.head()
18
19 df = df.drop(["Unnamed: 2", "Unnamed: 3", "Unnamed: 4"], axis=1)
20 df.columns = ["labels", "data"]
21 df["b_labels"] = df["labels"].map({"ham": 0, "spam": 1})
22 y = df["b_labels"].values
23
24 x_train, x_test, y_train, y_test = train_test_split(df["data"], y,
    test_size=0.33)
25
26 num_words = 2000
27 tokenizer = Tokenizer(num_words=num_words)

```

```

28 tokenizer.fit_on_texts(x_train)
29 sequences_train = tokenizer.texts_to_sequences(x_train)
30 sequences_test = tokenizer.texts_to_sequences(x_test)
31 word2index = tokenizer.word_index
32 V = len(word2index)
33 print("%s tokens" % V)
34
35 data_train = pad_sequences(sequences_train)
36 T = data_train.shape[1]
37 data_test = pad_sequences(sequences_test, maxlen=T)
38 print("data_train.shape = ", data_train.shape)
39 print("data_test.shape = ", data_test.shape)
40
41 D = 20
42 M = 5
43
44 i = Input(shape=(T,))
45 x = Embedding(V + 1, D)(i)
46 x = LSTM(M)(x)
47 x = Dense(1, activation="sigmoid")(x)
48
49 model = Model(i, x)
50
51 model.compile(
52     loss="binary_crossentropy",
53     optimizer="adam",
54     metrics=["accuracy"]
55 )
56
57 epochs = 5
58
59 r = model.fit(
60     data_train, y_train,
61     validation_data=(data_test, y_test),
62     epochs=epochs
63 )
64
65 plt.plot(r.history["loss"], label="loss")
66 plt.plot(r.history["val_loss"], label="val_loss")
67 plt.xlabel("Épocas")
68 plt.ylabel("loss")
69 plt.xticks(np.arange(0, epochs, 1), labels=range(1, epochs + 1))
70 plt.legend()
71 plt.show()

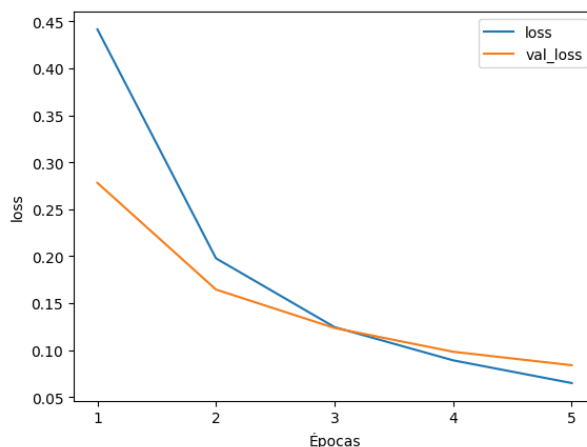
```

7242 tokens

data_train.shape = (3733, 175)

data_test.shape = (1839, 175)

Figura 8.5 – DETECTOR DE SPAM COM RNN: GRÁFICO DE PERDA

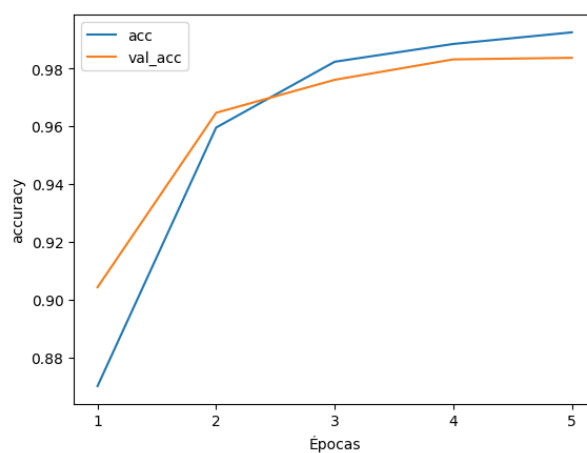


```

1 plt.plot(r.history["accuracy"], label="acc")
2 plt.plot(r.history["val_accuracy"], label="val_acc")
3 plt.xlabel("Épocas")
4 plt.ylabel("accuracy")
5 plt.xticks(np.arange(0, epochs, 1), labels=range(1, epochs + 1))
6 plt.legend()
7 plt.show()

```

Figura 8.6 – DETECTOR DE SPAM COM RNN: GRÁFICO DE ACURÁCIA



```

1 texto = "Is your car dirty? Discover our new product. Free for all.
   Click the link."
2 seq_texto = tokenizer.texts_to_sequences([texto])
3 data_texto = pad_sequences(seq_texto, maxlen=T)
4
5 pred = model.predict(data_texto)
6 print(pred)
7 print("SPAM" if pred >= 0.5 else "OK")

```

[[0.6281646]] SPAM

3 Gerador de Dígitos Fake

```

1 %pip install imageio
2 %pip install git+https://github.com/tensorflow/docs
3
4 import glob
5 import imageio
6 import numpy as np
7 import os
8 import PIL
9 import time
10 import tensorflow as tf
11 import matplotlib.pyplot as plt
12
13 from tensorflow.keras import layers
14 from IPython import display
15
16 (train_images, train_labels) , (_,_) = tf.keras.datasets.mnist.
    load_data()
17 train_images = train_images.reshape(train_images.shape[0], 28, 28, 1)
    .astype('float32')
18 train_images = (train_images - 127.5) / 127.5
19
20 BUFFER_SIZE = 60000
21 BATCH_SIZE = 256
22
23 train_dataset = tf.data.Dataset.from_tensor_slices(train_images).
    shuffle(BUFFER_SIZE).batch(BATCH_SIZE)
24
25 def make_generator_model():
26     model = tf.keras.Sequential()
27     model.add(layers.Dense(7*7*256, use_bias=False, input_shape=(100,))
        )
28
29     model.add(layers.BatchNormalization())
30     model.add(layers.LeakyReLU())
31
32     model.add(layers.Reshape((7, 7, 256)))
33     assert model.output_shape == (None, 7, 7, 256)
34
35     model.add(layers.Conv2DTranspose(128, (5, 5), strides=1, padding='
        same', use_bias=False))
36     assert model.output_shape == (None, 7, 7, 128)
37
38     model.add(layers.BatchNormalization())
39     model.add(layers.LeakyReLU())
40
41     model.add(layers.Conv2DTranspose(64, (5, 5), strides=(2, 2),
        padding='same', use_bias=False))
42     assert model.output_shape == (None, 14, 14, 64)
43
44     model.add(layers.BatchNormalization())
45     model.add(layers.LeakyReLU())

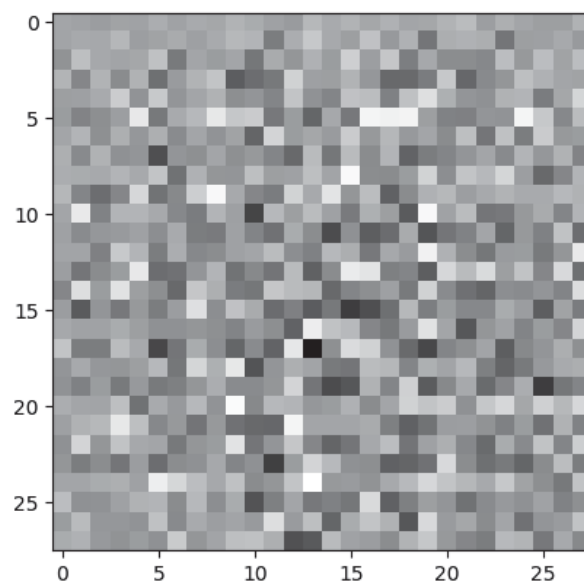
```

```

46 model.add(layers.Conv2DTranspose(1, (5, 5), strides=(2, 2), padding
47     ='same', use_bias=False, activation='tanh'))
48 assert model.output_shape == (None, 28, 28, 1)
49
50 return model
51
52 generator = make_generator_model()
53 noise = tf.random.normal([1, 100])
54 generated_image = generator(noise, training=False)
55 plt.imshow(generated_image[0, :, :, 0], cmap='gray')

```

Figura 8.7 – GERANDO DADOS COM GAN: GERAÇÃO DE IMAGEM RUIDOSA



```

1 def make_discriminator_model():
2     model = tf.keras.Sequential()
3     model.add(layers.Conv2D(64, (5, 5), strides=(2, 2), padding='same',
4         input_shape=[28, 28, 1]))
5     model.add(layers.LeakyReLU())
6     model.add(layers.Dropout(0.3))
7
8     model.add(layers.Conv2D(128, (5, 5), strides=(2, 2), padding='same'
9         ))
10    model.add(layers.LeakyReLU())
11    model.add(layers.Dropout(0.3))
12
13    model.add(layers.Flatten())
14    model.add(layers.Dense(1))
15
16    return model
17
18 discriminator = make_discriminator_model()
19 decision = discriminator(generated_image)

```

```

18
19 cross_entropy = tf.keras.losses.BinaryCrossentropy(from_logits=True)
20
21 def discriminator_loss(real_output, fake_output):
22     real_loss = cross_entropy(tf.ones_like(real_output), real_output)
23     fake_loss = cross_entropy(tf.zeros_like(fake_output), fake_output)
24     total_loss = real_loss + fake_loss
25     return total_loss
26
27 def generator_loss(fake_output):
28     return cross_entropy(tf.ones_like(fake_output), fake_output)
29
30 generator_optimizer = tf.keras.optimizers.Adam(1e-4)
31 discriminator_optimizer = tf.keras.optimizers.Adam(1e-4)
32
33 checkpoint_dir = './training_checkpoints'
34 checkpoint_prefix = os.path.join(checkpoint_dir, "ckpt")
35 checkpoint = tf.train.Checkpoint(generator_optimizer=
36     discriminator_optimizer=
37         discriminator_optimizer,
38         generator=generator,
39         discriminator=discriminator)
40
41 EPOCHS = 100
42 noise_dim = 100
43 num_examples_to_generate = 16
44
45 seed = tf.random.normal([num_examples_to_generate, noise_dim])
46
47 @tf.function
48 def train_step(images):
49     noise = tf.random.normal([BATCH_SIZE, noise_dim])
50
51     with tf.GradientTape() as gen_tape, tf.GradientTape() as disc_tape:
52         generated_images = generator(noise, training=True)
53
54         real_output = discriminator(images, training=True)
55         fake_output = discriminator(generated_images, training=True)
56
57         gen_loss = generator_loss(fake_output)
58         disc_loss = discriminator_loss(real_output, fake_output)
59
60         gradients_of_generator = gen_tape.gradient(gen_loss, generator.
61             trainable_variables)
62         gradients_of_discriminator = disc_tape.gradient(disc_loss,
63             discriminator.trainable_variables)
64
65         generator_optimizer.apply_gradients(zip(gradients_of_generator,
66             generator.trainable_variables))

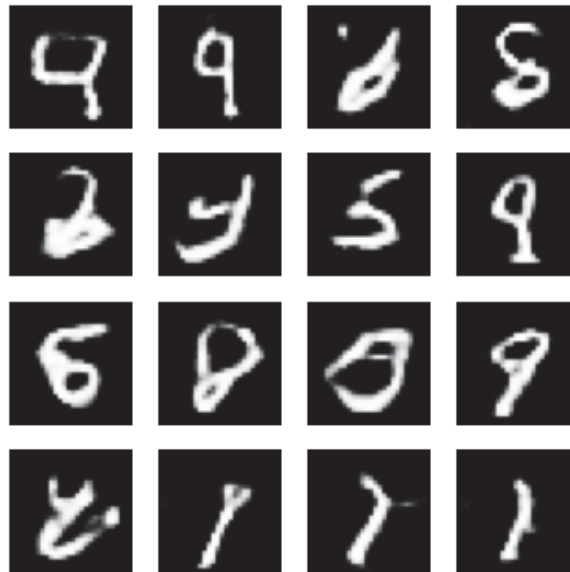
```

```

63     discriminator_optimizer.apply_gradients(zip(
        gradients_of_discriminator, discriminator.trainable_variables))
64
65 def train(dataset, epochs):
66     for epoch in range(epochs):
67         start = time.time()
68
69         for image_batch in dataset:
70             train_step(image_batch)
71
72         display.clear_output(wait=True)
73         generate_and_save_images(generator, epoch + 1, seed)
74
75         if (epoch + 1) % 15 == 0:
76             checkpoint.save(file_prefix = checkpoint_prefix)
77
78         print('Time for epoch {} is {} sec'.format(epoch + 1, time.time()
        - start))
79
80     display.clear_output(wait=True)
81     generate_and_save_images(generator, epochs, seed)
82
83 def generate_and_save_images(model, epoch, test_input):
84     predictions = model(test_input, training=False)
85     fig = plt.figure(figsize=(4, 4))
86
87     for i in range(predictions.shape[0]):
88         plt.subplot(4, 4, i + 1)
89         plt.imshow(predictions[i, :, :, 0] * 127.5 + 127.5, cmap='gray')
90         plt.axis('off')
91
92     plt.savefig('image_at_epoch_{:04d}.png'.format(epoch))
93     plt.show()
94
95 train(train_dataset, EPOCHS)
96 checkpoint.restore(tf.train.latest_checkpoint(checkpoint_dir))

```

Figura 8.8 – GERANDO DADOS COM GAN: CHECKPOINT DA GERAÇÃO DE DÍGITOS



```

1 def display_image(epoch_no):
2     return PIL.Image.open('image_at_epoch_{:04d}.png'.format(epoch_no))
3
4 display_image(EPOCHS)
5
6 anim_file = 'dcgan.gif'
7
8 with imageio.get_writer(anim_file, mode='I') as writer:
9     filenames = glob.glob('image*.png')
10    filenames = sorted(filenames)
11    for filename in filenames:
12        image = imageio.imread(filename)
13        writer.append_data(image)
14        image = imageio.imread(filename)
15        writer.append_data(image)
16
17 import tensorflow_docs.vis.embed as embed
18 embed.embed_file(anim_file)

```

4 Transformer

```

1 !pip uninstall tensorflow
2 !pip install tensorflow==2.15.0
3 !pip install tensorflow_datasets
4 !pip install -U tensorflow-text==2.15.0
5 import collections
6 import logging
7 import os
8 import pathlib
9 import re
10 import string
11 import sys

```

```

12 import time
13 import numpy as np
14 import matplotlib.pyplot as plt
15 import tensorflow_datasets as tfds
16 import tensorflow_text as text
17 import tensorflow as tf
18 logging.getLogger('tensorflow').setLevel(logging.ERROR)
19
20 examples, metadata = tfds.load('ted_hrlr_translate/pt_to_en',
21 with_info=True, as_supervised=True)
22
23 train_examples, val_examples = examples['train'], examples['
    validation']
24
25 for pt_examples, en_examples in train_examples.batch(3).take(1):
26     for pt in pt_examples.numpy():
27         print(pt.decode('utf-8'))
28
29     print()
30
31     for en in en_examples.numpy():
32         print(en.decode('utf-8'))

```

e quando melhoramos a procura , tiramos a única vantagem da impressão ,
que é a serendipidade .

mas e se estes fatores fossem ativos ?

mas eles não tinham a curiosidade de me testar .

and when you improve searchability , you actually take away the one
advantage of print , which is serendipity .

but what if it were active ?

but they did n't test for curiosity .

```

1 # Tokenização e Destokenização do texto
2 model_name = "ted_hrlr_translate_pt_en_converter"
3
4 tf.keras.utils.get_file(f"{model_name}.zip",
5 f"https://storage.googleapis.com/download.tensorflow.org/models/{
    model_name}.zip", cache_dir='.', cache_subdir='', extract=True)
6 # Tem 2 tokenizers: um pt outro em en
7 # tokenizers.en tokeniza e detokeniza
8 tokenizers = tf.saved_model.load(model_name)
9
10 # PIPELINE DE ENTRADA
11 # Codificar/tokenizar lotes de texto puro
12 def tokenize_pairs(pt, en):
13     pt = tokenizers.pt.tokenize(pt)
14     # Converte ragged (irregular, tam variável) para dense
15     # Faz padding com zeros.
16     pt = pt.to_tensor()
17
18     en = tokenizers.en.tokenize(en)
19     # ragged -> dense

```

```

20     en = en.to_tensor()
21
22     return pt, en
23
24 # Pipeline simples: processa, embaralha, agrupa os dados, prefetch
25 # Datasets de entrada terminam com prefetch
26 BUFFER_SIZE = 20000
27 BATCH_SIZE = 64
28
29 def make_batches(ds):
30     return (
31         ds
32         .cache()
33         .shuffle(BUFFER_SIZE)
34         .batch(BATCH_SIZE)
35         .map(tokenize_pairs, num_parallel_calls=tf.data.AUTOTUNE)
36         .prefetch(tf.data.AUTOTUNE))
37
38 train_batches = make_batches(train_examples)
39 val_batches = make_batches(val_examples)
40
41 # CODIFICAÇÃO POSICIONAL
42 def get_angles(pos, i, d_model):
43     angle_rates = 1 / np.power(10000, (2 * (i//2)) / np.float32(d_model
44         ))
45     return pos * angle_rates
46
47 def positional_encoding(position, d_model):
48     angle_rads = get_angles(np.arange(position)[:, np.newaxis], np.
49         arange(d_model)[np.newaxis, :], d_model)
50
51     # sin em índices pares no array; 2i
52     angle_rads[:, 0::2] = np.sin(angle_rads[:, 0::2])
53
54     # cos em índices ímpares no array; 2i+1
55     angle_rads[:, 1::2] = np.cos(angle_rads[:, 1::2])
56
57     # newaxis, aumenta a dimensão [] -> [ [] ]
58     pos_encoding = angle_rads[np.newaxis, ...]
59
60     return tf.cast(pos_encoding, dtype=tf.float32)
61
62 # CODIFICAÇÃO POSICIONAL
63 n, d = 2048, 512
64 pos_encoding = positional_encoding(n, d)
65 print(pos_encoding.shape)
66 pos_encoding = pos_encoding[0]
67
68 # Arrumar as dimensões
69 pos_encoding = tf.reshape(pos_encoding, (n, d//2, 2))
70 pos_encoding = tf.transpose(pos_encoding, (2, 1, 0))

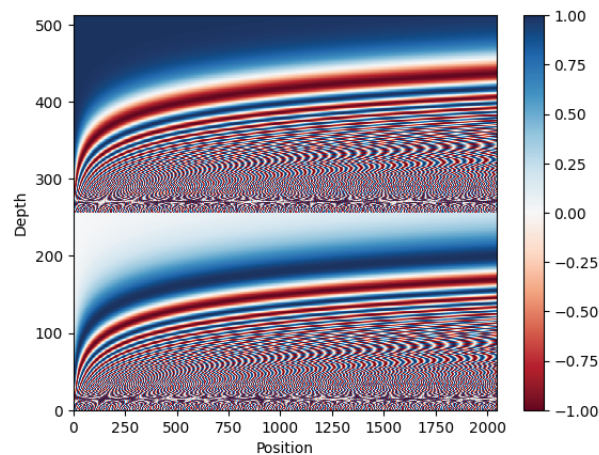
```

```

69 pos_encoding = tf.reshape(pos_encoding, (d, n))
70
71 plt.pcolormesh(pos_encoding, cmap='RdBu')
72 plt.ylabel('Depth')
73 plt.xlabel('Position')
74 plt.colorbar()
75 plt.show()

```

Figura 8.9 – TRANSFORMER: CODIFICAÇÃO POSICIONAL



```

1  # Cria uma máscara de 0 e 1, 0 para quando há valor e 1 quando não há
2  def create_padding_mask(seq):
3      seq = tf.cast(tf.math.equal(seq, 0), tf.float32)
4
5      # add extra dimensions to add the padding
6      # to the attention logits.
7      return seq[:, tf.newaxis, tf.newaxis, :] # (batch_size, 1, 1, seq_len)
8
9  # Máscara futura, usada no decoder
10 def create_look_ahead_mask(size):
11     # zera o triângulo inferior
12     mask = 1 - tf.linalg.band_part(tf.ones((size, size)), -1, 0)
13     return mask # (seq_len, seq_len)
14
15 # Função de Atenção
16 def scaled_dot_product_attention(q, k, v, mask):
17     # Q K^T
18     matmul_qk = tf.matmul(q, k, transpose_b=True) # (... , seq_len_q, seq_len_k)
19
20     # converte matmul_qk para float32
21     dk = tf.cast(tf.shape(k)[-1], tf.float32)
22
23     # divide por sqrt(d_k)
24     scaled_attention_logits = matmul_qk / tf.math.sqrt(dk)
25

```

```

26 # Soma a máscara, e os valores faltantes serão um número próximo a -
    inf
27 if mask is not None:
28     scaled_attention_logits += (mask * -1e9)
29
30 # softmax normaliza os dados, soman 1. // (... , seq_len_q,
    seq_len_k)
31 attention_weights = tf.nn.softmax(scaled_attention_logits, axis=-1)
32
33 output = tf.matmul(attention_weights, v) # (... , seq_len_q, depth_v
    )
34
35 return output, attention_weights

```

```

1 # Atenção Multi-cabeças
2 class MultiHeadAttention(tf.keras.layers.Layer):
3     def __init__(self, d_model, num_heads):
4         super(MultiHeadAttention, self).__init__()
5         self.num_heads = num_heads
6         self.d_model = d_model
7
8         assert d_model % self.num_heads == 0
9
10        self.depth = d_model // self.num_heads
11
12        self.wq = tf.keras.layers.Dense(d_model)
13        self.wk = tf.keras.layers.Dense(d_model)
14        self.wv = tf.keras.layers.Dense(d_model)
15
16        self.dense = tf.keras.layers.Dense(d_model)
17
18        def split_heads(self, x, batch_size):
19            """Separa a última dimensão em (num_heads, depth).
20            Transpõe o resultado para o shape (batch_size, num_heads, seq_len,
                depth)
21            """
22            x = tf.reshape(x, (batch_size, -1, self.num_heads, self.depth))
23
24            return tf.transpose(x, perm=[0, 2, 1, 3])
25
26        def call(self, v, k, q, mask):
27            batch_size = tf.shape(q)[0]
28
29            q = self.wq(q) # (batch_size, seq_len, d_model)
30            k = self.wk(k) # (batch_size, seq_len, d_model)
31            v = self.wv(v) # (batch_size, seq_len, d_model)
32
33            q = self.split_heads(q, batch_size) # (batch_size, num_heads,
                seq_len_q, depth)
34            k = self.split_heads(k, batch_size) # (batch_size, num_heads,
                seq_len_k, depth)

```

```

35     v = self.split_heads(v, batch_size) # (batch_size, num_heads,
      seq_len_v, depth)
36
37     # Calcula a atenção para cada cabeça (de forma matricial)
38     # scaled_attention.shape == (batch_size, num_heads, seq_len_q,
      depth)
39     # attention_weights.shape == (batch_size, num_heads, seq_len_q,
      seq_len_k)
40     scaled_attention, attention_weights = scaled_dot_product_attention
      (q, k, v, mask)
41
42     # Troca a dimensão 2 com 1, para acertar o num_heads
43     # (batch_size, seq_len_q, num_heads, depth)
44     scaled_attention = tf.transpose(scaled_attention, perm=[0, 2, 1,
      3])
45
46     # Concatena os valores em: (batch_size, seq_len_q, d_model)
47     concat_attention = tf.reshape(scaled_attention, (batch_size, -1,
      self.d_model))
48
49     output = self.dense(concat_attention) # (batch_size, seq_len_q,
      d_model)
50
51     return output, attention_weights
52
53 def point_wise_feed_forward_network(d_model, dff):
54     return tf.keras.Sequential([
55         tf.keras.layers.Dense(dff, activation='relu'), # (batch_size,
      seq_len, dff)
56         tf.keras.layers.Dense(d_model) # (batch_size, seq_len, d_model)
57     ])

```

```

1 class EncoderLayer(tf.keras.layers.Layer):
2     def __init__(self, d_model, num_heads, dff, rate=0.1):
3         super(EncoderLayer, self).__init__()
4
5         self.mha = MultiHeadAttention(d_model, num_heads)
6         self.ffn = point_wise_feed_forward_network(d_model, dff)
7
8         self.layernorm1 = tf.keras.layers.LayerNormalization(epsilon=1e-6)
9         self.layernorm2 = tf.keras.layers.LayerNormalization(epsilon=1e-6)
10
11         self.dropout1 = tf.keras.layers.Dropout(rate)
12         self.dropout2 = tf.keras.layers.Dropout(rate)
13
14     def call(self, x, training, mask):
15         attn_output, _ = self.mha(x, x, x, mask) # (batch_size,
      input_seq_len, d_model)
16         attn_output = self.dropout1(attn_output, training=training)
17         out1 = self.layernorm1(x + attn_output) # (batch_size,
      input_seq_len, d_model)

```

```

18
19     ffn_output = self.ffn(out1) # (batch_size, input_seq_len, d_model)
20     ffn_output = self.dropout2(ffn_output, training=training)
21     out2 = self.layer_norm2(out1 + ffn_output) # (batch_size,
        input_seq_len, d_model)
22
23     return out2

```

```

1 class DecoderLayer(tf.keras.layers.Layer):
2     def __init__(self, d_model, num_heads, dff, rate=0.1):
3         super(DecoderLayer, self).__init__()
4
5         self.mha1 = MultiHeadAttention(d_model, num_heads)
6         self.mha2 = MultiHeadAttention(d_model, num_heads)
7
8         self.ffn = point_wise_feed_forward_network(d_model, dff)
9
10        self.layer_norm1 = tf.keras.layers.LayerNormalization(epsilon=1e-6)
11        self.layer_norm2 = tf.keras.layers.LayerNormalization(epsilon=1e-6)
12        self.layer_norm3 = tf.keras.layers.LayerNormalization(epsilon=1e-6)
13
14        self.dropout1 = tf.keras.layers.Dropout(rate)
15        self.dropout2 = tf.keras.layers.Dropout(rate)
16        self.dropout3 = tf.keras.layers.Dropout(rate)
17
18        def call(self, x, enc_output, training, look_ahead_mask,
        padding_mask):
19            # enc_output.shape == (batch_size, input_seq_len, d_model)
20            # (batch_size, target_seq_len, d_model)
21            attn1, attn_weights_block1 = self.mha1(x, x, x, look_ahead_mask)
22            attn1 = self.dropout1(attn1, training=training)
23            out1 = self.layer_norm1(attn1 + x)
24
25            # (batch_size, target_seq_len, d_model)
26            attn2, attn_weights_block2 = self.mha2(enc_output, enc_output,
        out1, padding_mask)
27            attn2 = self.dropout2(attn2, training=training)
28            out2 = self.layer_norm2(attn2 + out1) # (batch_size, target_seq_len
        , d_model)
29
30            ffn_output = self.ffn(out2) # (batch_size, target_seq_len, d_model
        )
31            ffn_output = self.dropout3(ffn_output, training=training)
32            out3 = self.layer_norm3(ffn_output + out2) # (batch_size,
        target_seq_len, d_model)
33
34            return out3, attn_weights_block1, attn_weights_block2

```

```

1 class Encoder(tf.keras.layers.Layer):
2     def __init__(self, num_layers, d_model, num_heads, dff,
        input_vocab_size, maximum_position_encoding, rate=0.1):

```

```

3     super(Encoder, self).__init__()
4
5     self.d_model = d_model
6     self.num_layers = num_layers
7     self.embedding = tf.keras.layers.Embedding(input_vocab_size,
8         d_model)
9     self.pos_encoding = positional_encoding(maximum_position_encoding,
10         self.d_model)
11     self.enc_layers = [EncoderLayer(d_model, num_heads, dff, rate) for
12         _ in range(num_layers)]
13     self.dropout = tf.keras.layers.Dropout(rate)
14
15     def call(self, x, training, mask):
16         seq_len = tf.shape(x)[1]
17
18         # adding embedding and position encoding.
19         x = self.embedding(x) # (batch_size, input_seq_len, d_model)
20         x *= tf.math.sqrt(tf.cast(self.d_model, tf.float32))
21         x += self.pos_encoding[:, :seq_len, :]
22         x = self.dropout(x, training=training)
23
24         for i in range(self.num_layers):
25             x = self.enc_layers[i](x, training, mask)
26
27     return x # (batch_size, input_seq_len, d_model)

```

```

1 class Decoder(tf.keras.layers.Layer):
2     def __init__(self, num_layers, d_model, num_heads, dff,
3         target_vocab_size, maximum_position_encoding, rate=0.1):
4         super(Decoder, self).__init__()
5
6         self.d_model = d_model
7         self.num_layers = num_layers
8
9         self.embedding = tf.keras.layers.Embedding(target_vocab_size,
10             d_model)
11         self.pos_encoding = positional_encoding(maximum_position_encoding,
12             d_model)
13
14         self.dec_layers = [DecoderLayer(d_model, num_heads, dff, rate) for
15             _ in range(num_layers)]
16         self.dropout = tf.keras.layers.Dropout(rate)
17
18     def call(self, x, enc_output, training, look_ahead_mask,
19         padding_mask):
20         seq_len = tf.shape(x)[1]
21         attention_weights = {}
22
23         x = self.embedding(x) # (batch_size, target_seq_len, d_model)
24         x *= tf.math.sqrt(tf.cast(self.d_model, tf.float32))
25         x += self.pos_encoding[:, :seq_len, :]

```

```

21
22     x = self.dropout(x, training=training)
23
24     for i in range(self.num_layers):
25         x, block1, block2 = self.dec_layers[i](x, enc_output, training,
26         look_ahead_mask, padding_mask)
27
28         attention_weights[f'decoder_layer{i+1}_block1'] = block1
29         attention_weights[f'decoder_layer{i+1}_block2'] = block2
30
31     # x.shape == (batch_size, target_seq_len, d_model)
32     return x, attention_weights

```

```

1 class Transformer(tf.keras.Model):
2     def __init__(self, num_layers, d_model, num_heads, dff,
3         input_vocab_size, target_vocab_size, pe_input, pe_target, rate
4         =0.1):
5         super().__init__()
6         self.encoder = Encoder(num_layers, d_model, num_heads, dff,
7         input_vocab_size,
8         pe_input, rate)
9         self.decoder = Decoder(num_layers, d_model, num_heads, dff,
10         target_vocab_size,
11         pe_target, rate)
12         self.final_layer = tf.keras.layers.Dense(target_vocab_size)
13
14     def call(self, inputs, training):
15         # Keras models prefer if you pass all your inputs in the first
16         # argument
17         inp, tar = inputs
18
19         enc_padding_mask, look_ahead_mask, dec_padding_mask = self.
20         create_masks(inp, tar)
21
22         # (batch_size, inp_seq_len, d_model)
23         enc_output = self.encoder(inp, training, enc_padding_mask)
24
25         # dec_output.shape == (batch_size, tar_seq_len, d_model)
26         dec_output, attention_weights = self.decoder(tar, enc_output,
27         training, look_ahead_mask, dec_padding_mask)
28
29         # (batch_size, tar_seq_len, target_vocab_size)
30         final_output = self.final_layer(dec_output)
31
32     return final_output, attention_weights
33
34     def create_masks(self, inp, tar):
35         # Encoder padding mask
36         enc_padding_mask = create_padding_mask(inp)
37         # Used in the 2nd attention block in the decoder.
38         # This padding mask is used to mask the encoder outputs.

```

```

32 dec_padding_mask = create_padding_mask(inp)
33 # Used in the 1st attention block in the decoder.
34 # It is used to pad and mask future tokens in the input received
    by
35 # the decoder.
36 look_ahead_mask = create_look_ahead_mask(tf.shape(tar)[1])
37 dec_target_padding_mask = create_padding_mask(tar)
38 look_ahead_mask = tf.maximum(dec_target_padding_mask,
    look_ahead_mask)
39
40 return enc_padding_mask, look_ahead_mask, dec_padding_mask
41
42 # Hiperparâmetros
43 num_layers = 4
44 d_model = 128
45 dff = 512
46 num_heads = 8
47 dropout_rate = 0.1

```

```

1
2 class CustomSchedule(tf.keras.optimizers.schedules.
    LearningRateSchedule):
3     def __init__(self, d_model, warmup_steps=4000):
4         super(CustomSchedule, self).__init__()
5         self.d_model = d_model
6         self.d_model = tf.cast(self.d_model, tf.float32)
7         self.warmup_steps = warmup_steps
8
9     def __call__(self, step):
10        step = tf.cast(step, tf.float32) # Adicionado para evitar ERRO
11        arg1 = tf.math.rsqrt(step)
12        arg2 = step * (self.warmup_steps ** -1.5)
13        return tf.math.rsqrt(self.d_model) * tf.math.minimum(arg1, arg2)

```

```

1
2 learning_rate = CustomSchedule(d_model)
3 optimizer = tf.keras.optimizers.Adam(learning_rate, beta_1=0.9,
    beta_2=0.98, epsilon=1e-9)
4
5 loss_object = tf.keras.losses.SparseCategoricalCrossentropy(
    from_logits=True, reduction='none')
6
7 def loss_function(real, pred):
8     mask = tf.math.logical_not(tf.math.equal(real, 0))
9     loss_ = loss_object(real, pred)
10    mask = tf.cast(mask, dtype=loss_.dtype)
11    loss_ *= mask
12    return tf.reduce_sum(loss_) / tf.reduce_sum(mask)
13
14 def accuracy_function(real, pred):
15     accuracies = tf.equal(real, tf.argmax(pred, axis=2))

```

```

16 mask = tf.math.logical_not(tf.math.equal(real, 0))
17 accuracies = tf.math.logical_and(mask, accuracies)
18 accuracies = tf.cast(accuracies, dtype=tf.float32)
19 mask = tf.cast(mask, dtype=tf.float32)
20 return tf.reduce_sum(accuracies)/tf.reduce_sum(mask)
21
22 train_loss = tf.keras.metrics.Mean(name='train_loss')
23 train_accuracy = tf.keras.metrics.Mean(name='train_accuracy')
24
25 transformer = Transformer(
26     num_layers=num_layers,
27     d_model=d_model,
28     num_heads=num_heads,
29     dff=dff,
30     input_vocab_size=tokenizers.pt.get_vocab_size().numpy(),
31     target_vocab_size=tokenizers.en.get_vocab_size().numpy(),
32     pe_input=1000,
33     pe_target=1000,
34     rate=dropout_rate)
35
36 # Checkpoint
37 checkpoint_path = "./checkpoints/train"
38 ckpt = tf.train.Checkpoint(transformer=transformer, optimizer=
    optimizer)
39 ckpt_manager = tf.train.CheckpointManager(ckpt, checkpoint_path,
    max_to_keep=5)
40
41 # if a checkpoint exists, restore the latest checkpoint.
42 if ckpt_manager.latest_checkpoint:
43     ckpt.restore(ckpt_manager.latest_checkpoint)
44     print('Latest checkpoint restored!!')
45
46 EPOCHS = 20
47 train_step_signature = [
48     tf.TensorSpec(shape=(None, None), dtype=tf.int64),
49     tf.TensorSpec(shape=(None, None), dtype=tf.int64),
50 ]
51 @tf.function(input_signature=train_step_signature)
52 def train_step(inp, tar):
53     tar_inp = tar[:, :-1]
54     tar_real = tar[:, 1:]
55     with tf.GradientTape() as tape:
56         predictions, _ = transformer([inp, tar_inp], training = True)
57         loss = loss_function(tar_real, predictions)
58
59     gradients = tape.gradient(loss, transformer.trainable_variables)
60     optimizer.apply_gradients(zip(gradients, transformer.
        trainable_variables))
61
62     train_loss(loss)
63     train_accuracy(accuracy_function(tar_real, predictions))

```



```

27     predictions = predictions[:, -1:, :] # (batch_size, 1,
      vocab_size)
28     predicted_id = tf.argmax(predictions, axis=-1)
29     output_array = output_array.write(i+1, predicted_id[0])
30     if predicted_id == end:
31         break
32
33     output = tf.transpose(output_array.stack())
34     # output.shape (1, tokens)
35     text = tokenizers.en.detokenize(output)[0]
36     tokens = tokenizers.en.lookup(output)[0]
37     _, attention_weights = self.transformer([encoder_input, output
      [:-1]], training=False)
38
39     return text, tokens, attention_weights
40
41 translator = Translator(tokenizers, transformer)
42 sentence = "Eu li sobre triceratops na enciclopédia."
43 translated_text, translated_tokens, attention_weights = translator(tf
      .constant(sentence))
44
45 print(f'{"Prediction":15s}: {translated_text}')

```

Prediction : b'i read about trieps in the encycloody .'

APÊNDICE 9 – BIG DATA

A - ENUNCIADO

Enviar um arquivo PDF contendo uma descrição breve (2 páginas) sobre a implementação de uma aplicação ou estudo de caso envolvendo Big Data e suas ferramentas (NoSQL e NewSQL). Caracterize os dados e Vs envolvidos, além da modelagem necessária dependendo dos modelos de dados empregados.

B - RESOLUÇÃO

Estudo de Caso: A Aplicação de Big Data no Airbnb

Introdução

O Airbnb é uma plataforma online que conecta pessoas que desejam alugar suas casas com viajantes que buscam acomodação. Desde a sua fundação em 2008, a empresa cresceu exponencialmente, operando em mais de 190 países. Esse crescimento é sustentado pela capacidade do Airbnb de coletar, processar e analisar grandes volumes de dados, permitindo uma experiência personalizada para usuários e anfitriões. Neste estudo de caso, analisaremos como os dados do Airbnb podem ser classificados como Big Data, explorando os "Vs" do Big Data presentes em suas operações, e detalharemos as ferramentas utilizadas pela empresa no gerenciamento desses dados.

Classificação dos Dados do Airbnb como Big Data

Os dados gerados e utilizados pelo Airbnb apresentam características que os classificam como Big Data. Os cinco "Vs" do Big Data, Volume, Velocidade, Variedade, Veracidade e Valor, estão presentes de forma evidente nos dados da empresa.

1. **Volume:** O Airbnb lida com uma quantidade massiva de dados. São milhões de usuários, reservas, avaliações, mensagens e interações diariamente. Esse grande volume requer sistemas capazes de armazenar e processar dados em escala petabyte.
2. **Velocidade:** Os dados são gerados e precisam ser processados em tempo real ou quase real. Reservas, cancelamentos e atualizações de disponibilidade ocorrem constantemente, exigindo processamento rápido para manter a plataforma atualizada.
3. **Variedade:** Os dados do Airbnb são diversos, incluindo dados estruturados (informações de reservas, perfis de usuários), semiestruturados (logs de atividades) e não estruturados (comentários, imagens). Essa variedade demanda sistemas flexíveis para gerenciamento.
4. **Veracidade:** A qualidade e a confiabilidade dos dados são cruciais. O Airbnb precisa garantir que as informações sejam precisas para manter a confiança entre hóspedes e anfitriões.
5. **Valor:** A análise desses dados gera insights valiosos que permitem melhorar a experiência do usuário, otimizar preços e expandir os negócios. O valor extraído dos dados é fundamental para a estratégia da empresa.

Ferramentas Utilizadas

O Airbnb utiliza uma arquitetura sofisticada de Big Data, incorporando diversas ferramentas e tecnologias para atender às suas necessidades. Abaixo, detalhamos cada ferramenta e seu papel no processo de dados.

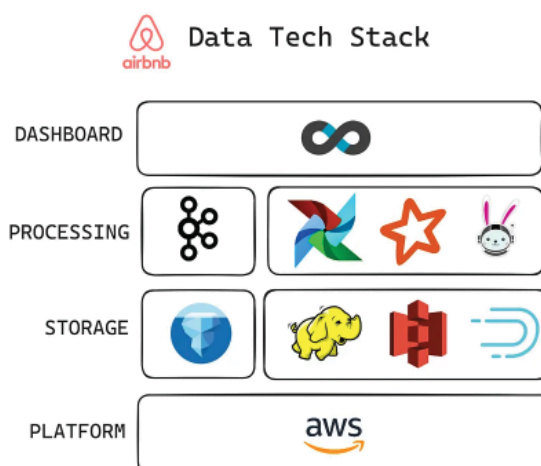
1. **Amazon Web Services (AWS):** S3 (Simple Storage Service): Utilizado para armazenamento de objetos, o S3 permite que o Airbnb armazene grandes volumes de dados de forma escalável e segura. E EC2 (Elastic Compute Cloud): Fornece capacidade computacional elástica, permitindo que a empresa escale seus recursos de processamento conforme a demanda.

2. Hadoop: O Hadoop é um framework de software que facilita o processamento distribuído de grandes conjuntos de dados. No Airbnb, é usado para armazenar e processar dados em clusters de servidores, permitindo a análise de grandes volumes de informação.
3. HBase: Um banco de dados NoSQL que funciona em cima do Hadoop, o HBase é usado para armazenamento em tempo real de dados grandes e dispersos. Ele permite acessos rápidos a dados específicos dentro de conjuntos massivos.
4. Hive: O Hive permite consultas tipo SQL em dados armazenados no Hadoop. Isso facilita para os analistas de dados executarem consultas complexas sem a necessidade de programação em baixo nível.
5. Presto: Presto é um motor de consultas distribuídas que permite ao Airbnb executar consultas SQL interativas e de baixa latência em grandes volumes de dados, espalhados por várias fontes de dados.
6. Apache Spark: O Spark é utilizado para processamento em lote e em tempo real. Com sua capacidade de processamento na memória, o Spark acelera o processamento de tarefas como análise de dados, aprendizado de máquina e processamento de gráficos.
7. Apache Airflow: Desenvolvido pelo próprio Airbnb, o Airflow é uma plataforma de orquestração de fluxos de trabalho que permite a automação de pipelines de dados. Ele gerencia a execução de tarefas, dependências e agendamento de processos.
8. Apache Superset: Outra criação do Airbnb, o Superset é uma plataforma de visualização de dados. Permite a criação de dashboards interativos e gráficos, auxiliando na interpretação e análise dos dados.
9. Druid: Druid é um sistema de banco de dados analítico projetado para consultas rápidas em grandes conjuntos de dados. No Airbnb, é usado para análise em tempo real e interativa, especialmente em dados de séries temporais.

Aplicação das Ferramentas

O processo de dados no Airbnb pode ser dividido em várias etapas: coleta, armazenamento, processamento, análise e visualização. A Figura abaixo mostra uma visão geral das ferramentas utilizadas em cada etapa, sendo estas especificadas nas seções seguintes.

Figura 9.1 – VISÃO GERAL DAS FERRAMENTAS UTILIZADAS PELO AIRBNB



1. **Coleta:** Dados são coletados de diversas fontes, incluindo interações de usuários, informações de reservas, avaliações, logs de atividades e dados de terceiros (como eventos locais e clima).
2. **Armazenamento:** S3 e Hadoop são utilizados para armazenar os dados coletados. O S3 armazena objetos de dados, enquanto o Hadoop permite armazenamento distribuído e escalável.
3. **Processamento:** Apache Spark processa dados em lote e em tempo real, permitindo análises complexas e treinamento de modelos de machine learning. Hive e Presto são usados para consultas e manipulação de dados, facilitando o acesso a informações específicas. HBase oferece acesso rápido a dados em tempo real, essencial para aplicações que requerem baixa latência.
4. **Análise:** Druid permite análises rápidas e interativas, especialmente útil para monitoramento em tempo real e análise de séries temporais. Modelos de machine learning são implementados para diversas finalidades, como previsão de demanda, definição de preços dinâmicos e recomendações personalizadas.
5. **Orquestração e Automação:** Apache Airflow gerencia e automatiza os fluxos de trabalho de dados, garantindo que os processos sejam executados corretamente e no tempo certo.
6. **Visualização:** Apache Superset é utilizado para criar dashboards e relatórios que auxiliam na tomada de decisões baseadas em dados.

Conclusão

O Airbnb é um exemplo notável de como Big Data pode ser utilizado para impulsionar negócios digitais. Através da coleta e análise de grandes volumes de dados, a empresa consegue oferecer experiências personalizadas, otimizar operações e inovar continuamente. Os cinco "Vs" do Big Data estão profundamente integrados nas operações do Airbnb, e o uso eficiente de ferramentas como AWS, Hadoop, Spark e outras permite que a empresa transforme dados brutos em insights valiosos. Como estudante de Big Data, observar o caso do Airbnb destaca a importância de uma arquitetura de dados robusta e da integração de diversas tecnologias para lidar com os desafios e oportunidades que o Big Data apresenta.

Referências

[https://www.linkedin.com/pulse/what-technology-stack-airbnb-john-doe?utm_source=share
utm_medium=member_android&utm_campaign=share_via](https://www.linkedin.com/pulse/what-technology-stack-airbnb-john-doe?utm_source=share&utm_medium=member_android&utm_campaign=share_via)
<https://www.junaideffendi.com/p/airbnb-data-tech-stack>
<https://vutr.substack.com/p/groupby-40-data-infrastructure-at>
<https://medium.com/airbnb-engineering/data-infrastructure-at-airbnb-8adfb34f169c>

APÊNDICE 10 – VISÃO COMPUTACIONAL

A - ENUNCIADO

Extração de Características

Os bancos de imagens fornecidos são conjuntos de imagens de 250x250 pixels de imuno-histoquímica (biópsia) de câncer de mama. No total são 4 classes (0, 1+, 2+ e 3+) que estão divididas em diretórios. O objetivo é classificar as imagens nas categorias correspondentes. Uma base de imagens será utilizada para o treinamento e outra para o teste do treino.

As imagens fornecidas são recortes de uma imagem maior do tipo WSI (Whole Slide Imaging) disponibilizada pela Universidade de Warwick (link). A nomenclatura das imagens segue o padrão XX_HER_YYYY.png, onde XX é o número do paciente e YYYY é o número da imagem recortada. Separe a base de treino em 80% para treino e 20% para validação. **Separe por pacientes (XX), não utilize a separação randômica! Pois, imagens do mesmo paciente não podem estar na base de treino e de validação, pois isso pode gerar um viés.** No caso da CNN VGG16 remova a última camada de classificação e armazene os valores da penúltima camada como um vetor de características. Após o treinamento, os modelos treinados devem ser validados na base de teste.

Tarefas:

- Carregue a base de dados de Treino.
- Crie partições contendo 80% para treino e 20% para validação (atenção aos pacientes).
- Extraia características utilizando LBP e a CNN VGG16 (gerando um csv para cada extrator).
- Treine modelos Random Forest, SVM e RNA para predição dos dados extraídos.
- Carregue a base de Teste e execute a tarefa 3 nesta base.
- Aplique os modelos treinados nos dados de treino
- Calcule as métricas de Sensibilidade, Especificidade e F1-Score com base em suas matrizes de confusão.
- Indique qual modelo dá o melhor o resultado e a métrica utilizada

Redes Neurais

Utilize as duas bases do exercício anterior para treinar as Redes Neurais Convolucionais VGG16 e a Resnet50. Utilize os pesos pré-treinados (Transfer Learning), refaça as camadas Fully Connected para o problema de 4 classes. Compare os treinos de 15 épocas com e sem Data Augmentation. Tanto a VGG16 quanto a Resnet50 têm como camada de entrada uma imagem 224x224x3, ou seja, uma imagem de 224x224 pixels coloridos (3 canais de cores). Portanto, será necessário fazer uma transformação de 250x250x3 para 224x224x3. Ao fazer o Data Augmentation cuidado para não alterar demais as cores das imagens e atrapalhar na classificação.

Tarefas:

- Utilize a base de dados de Treino já separadas em treino e validação do exercício anterior
- Treine modelos VGG16 e Resnet50 adaptadas com e sem Data Augmentation
- Aplique os modelos treinados nas imagens da base de Teste
- Calcule as métricas de Sensibilidade, Especificidade e F1-Score com base em suas matrizes de confusão.

e) Indique qual modelo dá o melhor o resultado e a métrica utilizada

B - RESOLUÇÃO

```

1
2 import os
3 import cv2
4 import shutil
5 import numpy as np
6 import pandas as pd
7 import matplotlib.pyplot as plt
8
9 from tqdm import tqdm
10 from pathlib import Path
11
12 from sklearn.svm import SVC
13 from sklearn.metrics import recall_score
14 from sklearn.preprocessing import LabelEncoder
15 from sklearn.ensemble import RandomForestClassifier
16 from sklearn.model_selection import train_test_split
17 from sklearn.metrics import classification_report, confusion_matrix
18
19 from tensorflow.keras.layers import Dense
20 from tensorflow.keras.models import Model
21 from tensorflow.keras.models import Sequential
22 from tensorflow.keras.applications import VGG16
23 from tensorflow.keras.preprocessing import image
24 from tensorflow.keras.applications.vgg16 import preprocess_input
25
26 !unzip Train_Warwick.zip
27
28 # Diretório das amostras de treino
29 images_dir = Path('/content/Train_4cls_amostra')
30
31 # Diretórios de destino para treino, validação e teste
32 train_dir = images_dir.parent / 'train'
33 val_dir = images_dir.parent / 'val'
34
35 # Classes disponíveis
36 classes = ['0', '1', '2', '3']
37
38 patient_dict = {}
39
40 # Percorrer cada classe e organizar pacientes
41 for cls in classes:
42     class_dir = images_dir / cls
43     for img_name in os.listdir(class_dir):
44         if img_name.endswith('.png'):
45             # Extrair o ID do paciente do nome do arquivo (XX_HER_YYYY.
46             # png)
47             patient_id = img_name.split('_')[0]

```

```

47         img_path = class_dir / img_name
48         if cls not in patient_dict:
49             patient_dict[cls] = {}
50         if patient_id not in patient_dict[cls]:
51             patient_dict[cls][patient_id] = []
52         patient_dict[cls][patient_id].append((img_path, cls))
53
54     # Função para dividir pacientes em 80% treino e 20% validação
55     def split_patients(patients):
56         patient_ids = list(patients.keys())
57         train_size = int(0.8 * len(patient_ids))
58         train_patient_ids = patient_ids[:train_size]
59         val_patient_ids = patient_ids[train_size:]
60         return train_patient_ids, val_patient_ids
61
62     # Função para mover as imagens para os diretórios de treino e validação
63     def move_images(patient_ids, patient_dict, target_dir):
64         for pid in patient_ids:
65             for img_path, cls in patient_dict[pid]:
66                 # Diretório da classe no conjunto alvo
67                 class_dir = target_dir / cls
68                 class_dir.mkdir(parents=True, exist_ok=True)
69                 # Caminho de destino da imagem
70                 dest_path = class_dir / img_path.name
71                 # Mover a imagem
72                 shutil.move(str(img_path), str(dest_path))
73
74     # Para cada classe, dividir pacientes e mover as imagens
75     for cls in classes:
76         print(f"Processando classe {cls}...")
77
78         # Obter os pacientes dessa classe
79         patients_in_class = patient_dict[cls]
80
81         # Dividir em treino e validação
82         train_patient_ids, val_patient_ids = split_patients(
            patients_in_class)
83
84         print('Número total de pacientes:', len(train_patient_ids) + len(
            val_patient_ids))
85         print('Número de pacientes para treino:', len(train_patient_ids))
86         print('Número de pacientes para validação:', len(val_patient_ids))
87
88         # Mover as imagens de treino
89         move_images(train_patient_ids, patients_in_class, train_dir)
90
91         # Mover as imagens de validação
92         move_images(val_patient_ids, patients_in_class, val_dir)
93
94     print("Divisão de treino e validação concluída.")

```

```

95
96 # Processando classe 0...
97 # Número total de pacientes: 5
98 # Número de pacientes para treino: 4
99 # Número de pacientes para validação: 1
100 # Processando classe 1...
101 # Número total de pacientes: 5
102 # Número de pacientes para treino: 4
103 # Número de pacientes para validação: 1
104 # Processando classe 2...
105 # Número total de pacientes: 5
106 # Número de pacientes para treino: 4
107 # Número de pacientes para validação: 1
108 # Processando classe 3...
109 # Número total de pacientes: 5
110 # Número de pacientes para treino: 4
111 # Número de pacientes para validação: 1
112 # Divisão de treino e validação concluída.
113
114 def get_pixel(img, center, x, y):
115
116     new_value = 0
117
118     try:
119         # If local neighbourhood pixel
120         # value is greater than or equal
121         # to center pixel values then
122         # set it to 1
123         if img[x][y] >= center:
124             new_value = 1
125
126     except:
127         # Exception is required when
128         # neighbourhood value of a center
129         # pixel value is null i.e. values
130         # present at boundaries.
131         pass
132
133     return new_value
134
135
136 # Function for calculating LBP
137 def lbp_calculated_pixel(img, x, y, radius):
138
139     center = img[x][y]
140
141     val_ar = []
142
143     # top_left
144     val_ar.append(get_pixel(img, center, x-radius, y-radius))
145

```

```

146     # top
147     val_ar.append(get_pixel(img, center, x-radius, y))
148
149     # top_right
150     val_ar.append(get_pixel(img, center, x-radius, y + radius))
151
152     # right
153     val_ar.append(get_pixel(img, center, x, y + radius))
154
155     # bottom_right
156     val_ar.append(get_pixel(img, center, x + radius, y + radius))
157
158     # bottom
159     val_ar.append(get_pixel(img, center, x + radius, y))
160
161     # bottom_left
162     val_ar.append(get_pixel(img, center, x + radius, y-radius))
163
164     # left
165     val_ar.append(get_pixel(img, center, x, y-radius))
166
167     # Now, we need to convert binary
168     # values to decimal
169     power_val = [1, 2, 4, 8, 16, 32, 64, 128]
170
171     val = 0
172
173     for i in range(len(val_ar)):
174         val += val_ar[i] * power_val[i]
175
176     return val
177
178 # Função para calcular o LBP
179 def extract_lbp_image(image):
180     img_gray = cv2.cvtColor(image,
181                             cv2.COLOR_BGR2GRAY)
182     img_gray = cv2.bilateralFilter(img_gray, 7, 75, 75)
183
184     height, width, _ = image.shape
185
186     img_lbp = np.zeros((height, width),
187                       np.uint8)
188
189     for i in range(0, height):
190         for j in range(0, width):
191             img_lbp[i, j] = lbp_calculated_pixel(img_gray, i, j, 1)
192
193     # Calcular o histograma
194     h = cv2.calcHist([img_lbp], [0], None, [256], [0, 256])
195     hist,bins = np.histogram(h.flatten(),256,[0,256])

```

```

196     h_normalized = 100*((hist - hist.min()) / (hist.max() - hist.min())
197         )
198     return img_lbp, h_normalized
199
200 def extract_lbp_from_directory(directory, classes):
201     features = []
202     labels = []
203
204     print(f"Extraíndo características LBP das imagens no diretório: {
205         directory}")
206
207     # Iterar sobre cada classe
208     for cls in classes:
209         class_dir = os.path.join(directory, cls)
210         image_names = os.listdir(class_dir)
211         for img_name in tqdm(image_names, desc=f"Classe {cls}"):
212             img_path = os.path.join(class_dir, img_name)
213             # Ler a imagem
214             image = cv2.imread(img_path, 1)
215             # Verificar se a imagem tem 4 canais
216             if len(image.shape) == 3 and image.shape[2] == 4:
217                 # Remover o canal alfa
218                 image = image[:, :, :3]
219                 # Extrair características LBP
220                 img_lbp, h_normalized = extract_lbp_image(image)
221                 features.append(h_normalized)
222                 # Armazenar o label
223                 labels.append(cls)
224
225     return features, labels
226
227 # Extrair características do conjunto de treino
228 train_features, train_labels = extract_lbp_from_directory(train_dir,
229     classes)
230
231 # Extrair características do conjunto de validação
232 val_features, val_labels = extract_lbp_from_directory(val_dir,
233     classes)
234
235 # Extrair características LBP das imagens no diretório: /content/
236 val
237
238 # Converter as listas em arrays numpy
239 train_features = np.array(train_features)
240 train_labels = np.array(train_labels)
241 val_features = np.array(val_features)
242 val_labels = np.array(val_labels)
243
244 # Criar DataFrames
245 train_df = pd.DataFrame(train_features)

```

```

242 train_df['label'] = train_labels
243
244 val_df = pd.DataFrame(val_features)
245 val_df['label'] = val_labels
246
247 # Salvar em CSV
248 train_df.to_csv('train_lbp_features.csv', index=False)
249 val_df.to_csv('val_lbp_features.csv', index=False)
250
251 print("Características LBP salvas nos arquivos 'train_lbp_features.
    csv' e 'val_lbp_features.csv'")
252
253 # Características LBP salvas nos arquivos 'train_lbp_features.csv' e
    'val_lbp_features.csv'
254
255 # Carregar o modelo VGG16 pré-treinado, sem a camada de classificação
    final
256 # pooling='avg' para obter um vetor de características de saída
257 base_model = VGG16(weights='imagenet', include_top=False, pooling='
    avg')
258
259 # Função para carregar uma imagem, redimensioná-la e pré-processá-la
260 def process_image(img_path):
261     img = image.load_img(img_path, target_size=(224, 224))
262     img_array = image.img_to_array(img)
263     img_array = np.expand_dims(img_array, axis=0)
264     img_array = preprocess_input(img_array) # Pré-processamento VGG16
265     return img_array
266
267 # Função para extrair as características usando o modelo VGG16
268 def extract_vgg16_features(img_path, model):
269     img_array = process_image(img_path)
270     features = model.predict(img_array)
271     return features.flatten() # Converter para vetor de caracterí
        sticas
272
273 # Extrair características de um diretório de imagens
274 def extract_features_from_directory(directory, model):
275     features = []
276     labels = []
277
278     for class_label in os.listdir(directory):
279         class_dir = os.path.join(directory, class_label)
280
281         if os.path.isdir(class_dir):
282             for img_name in tqdm(os.listdir(class_dir), desc=f"Classe {
                class_label}"):
283                 img_path = os.path.join(class_dir, img_name)
284
285                 # Extrair características
286                 img_features = extract_vgg16_features(img_path, model)

```

```

287
288         # Armazenar as características e o rótulo (classe)
289         features.append(img_features)
290         labels.append(class_label)
291
292     return np.array(features), np.array(labels)
293
294 # Extração das características das imagens de treino usando o modelo
    VGG16
295 print("Extraindo características do conjunto de treino usando o
    modelo VGG16")
296 train_features, train_labels = extract_features_from_directory(
    train_dir, base_model)
297
298 # Extração das características das imagens de validação usando o
    modelo VGG16
299 print("Extraindo características do conjunto de validação usando o
    modelo VGG16")
300 val_features, val_labels = extract_features_from_directory(val_dir,
    base_model)
301
302
303 # Converter para DataFrame e salvar como CSV
304 train_df = pd.DataFrame(train_features)
305 train_df['label'] = train_labels
306 train_df.to_csv('train_vgg16_features.csv', index=False)
307
308 val_df = pd.DataFrame(val_features)
309 val_df['label'] = val_labels
310 val_df.to_csv('val_vgg16_features.csv', index=False)
311
312 print("Características VGG16 salvas nos arquivos '
    train_vgg16_features.csv' e 'val_vgg16_features.csv'")
313
314 # Carregar características LBP e CNN VGG16
315 train_lbp = pd.read_csv('train_lbp_features.csv')
316 train_vgg16 = pd.read_csv('train_vgg16_features.csv')
317
318 # Separar características e rótulos (labels)
319 X_train_lbp = train_lbp.drop(columns=['label'])
320 y_train_lbp = train_lbp['label']
321
322 X_train_vgg16 = train_vgg16.drop(columns=['label'])
323 y_train_vgg16 = train_vgg16['label']
324
325 # Codificar as labels para valores numéricos
326 label_encoder = LabelEncoder()
327 y_train_lbp = label_encoder.fit_transform(y_train_lbp)
328 y_train_vgg16 = label_encoder.fit_transform(y_train_vgg16)
329
330 # Treinando o modelo SVM com características LBP

```

```

331 svm_lbp = SVC(kernel='linear', random_state=42)
332 svm_lbp.fit(X_train_lbp, y_train_lbp)
333
334 # Treinando o modelo SVM com características VGG16
335 svm_vgg16 = SVC(kernel='linear', random_state=42)
336 svm_vgg16.fit(X_train_vgg16, y_train_vgg16)
337
338 # Função para criar o modelo RNA
339 def create_rna_model(input_dim):
340     model = Sequential()
341     model.add(Dense(128, activation='relu', input_dim=input_dim))
342     model.add(Dense(64, activation='relu'))
343     model.add(Dense(32, activation='relu'))
344     model.add(Dense(len(label_encoder.classes_), activation='softmax'))
345     # Saída com número de classes
346     model.compile(optimizer='adam', loss='
        sparse_categorical_crossentropy', metrics=['accuracy'])
347     return model
348
349 # Treinando a RNA com características LBP
350 rna_lbp = create_rna_model(X_train_lbp.shape[1])
351 rna_lbp.fit(X_train_lbp, y_train_lbp, epochs=10, batch_size=32,
352             validation_split=0.2)
353
354 # Treinando a RNA com características VGG16
355 rna_vgg16 = create_rna_model(X_train_vgg16.shape[1])
356 rna_vgg16.fit(X_train_vgg16, y_train_vgg16, epochs=10, batch_size=32,
357               validation_split=0.2)
358
359 # Avaliando Random Forest (LBP)
360 y_pred_rf_lbp = rf_lbp.predict(X_train_lbp)
361 print("Relatório de Classificação (Random Forest - LBP):")
362 print(classification_report(y_train_lbp, y_pred_rf_lbp))
363
364 # Avaliando SVM (LBP)
365 y_pred_svm_lbp = svm_lbp.predict(X_train_lbp)
366 print("Relatório de Classificação (SVM - LBP):")
367 print(classification_report(y_train_lbp, y_pred_svm_lbp))
368
369 # Avaliando RNA (LBP)
370 y_pred_rna_lbp = rna_lbp.predict(X_train_lbp)
371 y_pred_rna_lbp = np.argmax(y_pred_rna_lbp, axis=1)
372 print("Relatório de Classificação (RNA - LBP):")
373 print(classification_report(y_train_lbp, y_pred_rna_lbp))
374
375 # Avaliando Random Forest (VGG16)
376 y_pred_rf_vgg16 = rf_vgg16.predict(X_train_vgg16)
377 print("Relatório de Classificação (Random Forest - VGG16):")
378 print(classification_report(y_train_vgg16, y_pred_rf_vgg16))
379
380 # Avaliando SVM (VGG16)

```

```

378 y_pred_svm_vgg16 = svm_vgg16.predict(X_train_vgg16)
379 print("Relatório de Classificação (SVM - VGG16):")
380 print(classification_report(y_train_vgg16, y_pred_svm_vgg16))
381
382 # Avaliando RNA (VGG16)
383 y_pred_rna_vgg16 = rna_vgg16.predict(X_train_vgg16)
384 y_pred_rna_vgg16 = np.argmax(y_pred_rna_vgg16, axis=1)
385 print("Relatório de Classificação (RNA - VGG16):")
386 print(classification_report(y_train_vgg16, y_pred_rna_vgg16))
387
388 #Relatório de Classificação (Random Forest - LBP):
389 #           precision recall f1-score support
390 #
391 #         0         1.00      1.00      1.00        116
392 #         1         1.00      1.00      1.00        117
393 #         2         1.00      1.00      1.00        120
394 #         3         1.00      1.00      1.00        120
395
396 #   accuracy                1.00      473
397 #  macro avg       1.00      1.00      1.00      473
398 #weighted avg       1.00      1.00      1.00      473
399
400 #Relatório de Classificação (SVM - LBP):
401 #           precision recall f1-score support
402 #
403 #         0         1.00      1.00      1.00        116
404 #         1         1.00      1.00      1.00        117
405 #         2         1.00      1.00      1.00        120
406 #         3         1.00      1.00      1.00        120
407
408 #   accuracy                1.00      473
409 #  macro avg       1.00      1.00      1.00      473
410 #weighted avg       1.00      1.00      1.00      473
411
412 #Relatório de Classificação (RNA - LBP):
413 #           precision recall f1-score support
414 #
415 #         0         0.76      0.73      0.75        116
416 #         1         0.71      0.89      0.79        117
417 #         2         0.43      0.74      0.54        120
418 #         3         1.00      0.07      0.12        120
419 #
420 #   accuracy                0.60      473
421 #  macro avg       0.73      0.61      0.55      473
422 # weighted avg       0.73      0.60      0.55      473
423
424 #Relatório de Classificação (Random Forest - VGG16):
425 #           precision recall f1-score support
426 #
427 #         0         1.00      1.00      1.00        116
428 #         1         1.00      1.00      1.00        117

```

```

429 #          2          1.00      1.00      1.00      120
430 #          3          1.00      1.00      1.00      120
431
432 # accuracy                      1.00      473
433 # macro avg      1.00      1.00      1.00      473
434 #weighted avg  1.00      1.00      1.00      473
435
436 # Relatório de Classificação (SVM - VGG16):
437 #           precision recall f1-score support
438
439 #          0          1.00      1.00      1.00      116
440 #          1          1.00      1.00      1.00      117
441 #          2          1.00      1.00      1.00      120
442 #          3          1.00      1.00      1.00      120
443
444 # accuracy                      1.00      473
445 # macro avg      1.00      1.00      1.00      473
446 #weighted avg  1.00      1.00      1.00      473
447
448 #Relatório de Classificação (RNA - VGG16):
449 #           precision recall f1-score support
450
451 #          0          0.98      1.00      0.99      116
452 #          1          0.92      1.00      0.96      117
453 #          2          1.00      0.85      0.92      120
454 #          3          0.95      1.00      0.98      120
455
456 # accuracy                      0.96      473
457 # macro avg      0.96      0.96      0.96      473
458 #weighted avg  0.96      0.96      0.96      473
459
460 # !unzip Test_Warwick.zip
461
462 # Diretório das amostras de teste
463 test_dir = Path('/content/Test_4cl_amostra')
464
465 # Extrair características do conjunto de teste
466 test_features, test_labels = extract_lbp_from_directory(test_dir,
467                                                         classes)
468
469 # Converter as listas em arrays numpy
470 test_features = np.array(test_features)
471 test_labels = np.array(test_labels)
472
473 # Criar DataFrame
474 test_df = pd.DataFrame(test_features)
475 test_df['label'] = test_labels
476
477 # Salvar em CSV
478 test_df.to_csv('test_lbp_features.csv', index=False)

```

```

479 print("Características LBP salvas no arquivo 'test_lbp_features.csv' "
      )
480
481 # Extração das características das imagens de teste
482 print("Extraindo características do conjunto de teste...")
483 test_features, test_labels = extract_features_from_directory(test_dir
      , base_model)
484
485 # Converter para DataFrame e salvar como CSV
486 test_df = pd.DataFrame(test_features)
487 test_df['label'] = test_labels
488 test_df.to_csv('test_vgg16_features.csv', index=False)
489
490 print("Características VGG16 salvas nos arquivos 'test_vgg16_features
      .csv' ")
491
492 # Carregar os dados de teste
493 test_lbp = pd.read_csv('test_lbp_features.csv')
494 test_vgg16 = pd.read_csv('test_vgg16_features.csv')
495
496 # Separar as características e os rótulos (labels)
497 X_test_lbp = test_lbp.drop(columns=['label'])
498 y_test_lbp = test_lbp['label']
499
500 X_test_vgg16 = test_vgg16.drop(columns=['label'])
501 y_test_vgg16 = test_vgg16['label']
502
503 # Codificar os rótulos de teste (mesma codificação usada para os
      dados de treino)
504 label_encoder = LabelEncoder()
505 y_test_lbp = label_encoder.fit_transform(y_test_lbp)
506
507 label_encoder = LabelEncoder()
508 y_test_vgg16 = label_encoder.fit_transform(y_test_vgg16)
509
510 # Fazer previsões com o modelo Random Forest
511 y_pred_rf_lbp = rf_lbp.predict(X_test_lbp)
512
513 y_pred_rf_vgg16 = rf_vgg16.predict(X_test_vgg16)
514
515 # Fazer previsões com o modelo SVM
516 y_pred_svm_lbp = svm_lbp.predict(X_test_lbp)
517
518 y_pred_svm_vgg16 = svm_vgg16.predict(X_test_vgg16)
519
520 # Fazer previsões com a RNA
521 y_pred_rna_lbp = rna_lbp.predict(X_test_lbp)
522 y_pred_rna_lbp = np.argmax(y_pred_rna_lbp, axis=1)
523
524 y_pred_rna_vgg16 = rna_vgg16.predict(X_test_vgg16)
525 y_pred_rna_vgg16 = np.argmax(y_pred_rna_vgg16, axis=1)

```

```

526
527 # Calcular Especificidade a partir da Matriz de Confusão
528 def calculate_specificity(conf_matrix):
529     TN = conf_matrix[0, 0]
530     FP = conf_matrix[0, 1]
531     specificity = TN / (TN + FP)
532     return specificity
533
534 print("Random Forest\n")
535
536 # Avaliar o desempenho do Random Forest (LBP)
537 print("LBP")
538 print(classification_report(y_test_lbp, y_pred_rf_lbp))
539 conf_matrix_rf_lbp = confusion_matrix(y_test_lbp, y_pred_rf_lbp)
540 specificity_rf_lbp = calculate_specificity(conf_matrix_rf_lbp)
541 print(f"Especificidade: {specificity_rf_lbp}")
542
543 # Avaliar o desempenho do Random Forest (VGG16)
544 print("\n\nVGG16")
545 print(classification_report(y_test_vgg16, y_pred_rf_vgg16))
546 conf_matrix_rf_vgg16 = confusion_matrix(y_test_vgg16, y_pred_rf_vgg16
    )
547 specificity_rf_vgg16 = calculate_specificity(conf_matrix_rf_vgg16)
548 print(f"Especificidade: {specificity_rf_vgg16}")
549
550
551 #Random Forest
552
553 #LBP
554 #           precision recall f1-score support
555 #
556 #      0      0.38    0.50    0.43    101
557 #      1      0.29    0.27    0.28     90
558 #      2      0.43    0.29    0.35     90
559 #      3      0.77    0.82    0.80     90
560
561 #   accuracy                0.47    371
562 #  macro avg    0.47    0.47    0.46    371
563 #weighted avg    0.47    0.47    0.46    371
564
565 #Especificidade: 0.6493506493506493
566
567
568 #VGG16
569 #           precision recall f1-score support
570 #
571 #      0      0.80    0.99    0.88    101
572 #      1      0.58    0.17    0.26     90
573 #      2      0.59    0.83    0.69     90
574 #      3      0.95    0.98    0.96     90
575

```

```

576 # accuracy 0.75 371
577 # macro avg 0.73 0.74 0.70 371
578 #weighted avg 0.73 0.75 0.70 371
579
580 #Especificidade: 0.9900990099009901
581
582 print("SVM\n")
583
584 # Avaliar o desempenho do SVM (LBP)
585 print("LBP")
586 print(classification_report(y_test_lbp, y_pred_svm_lbp))
587 conf_matrix_svm_lbp = confusion_matrix(y_test_lbp, y_pred_svm_lbp)
588 specificity_svm_lbp = calculate_specificity(conf_matrix_svm_lbp)
589 print(f"Especificidade: {specificity_svm_lbp}")
590
591 # Avaliar o desempenho do SVM (LBP)
592 print("\n\nVGG16")
593 print(classification_report(y_test_vgg16, y_pred_svm_vgg16))
594 conf_matrix_svm_vgg16 = confusion_matrix(y_test_vgg16,
595     y_pred_svm_vgg16)
596 specificity_svm_vgg16 = calculate_specificity(conf_matrix_svm_vgg16)
597 print(f"Especificidade: {specificity_svm_vgg16}")
598
599
600 #LBP
601 # precision recall f1-score support
602
603 # 0 0.38 0.50 0.43 101
604 # 1 0.35 0.32 0.34 90
605 # 2 0.50 0.43 0.46 90
606 # 3 0.74 0.63 0.68 90
607
608 # accuracy 0.47 371
609 # macro avg 0.49 0.47 0.48 371
610 #weighted avg 0.49 0.47 0.48 371
611
612 #Especificidade: 0.6538461538461539
613
614 #VGG16
615 # precision recall f1-score support
616
617 # 0 0.85 0.92 0.88 101
618 # 1 0.77 0.30 0.43 90
619 # 2 0.64 0.96 0.77 90
620 # 3 0.98 1.00 0.99 90
621
622 # accuracy 0.80 371
623 # macro avg 0.81 0.79 0.77 371
624 #weighted avg 0.81 0.80 0.77 371
625

```

```

626 #Especificidade: 0.9393939393939394
627
628 print("RNA\n")
629
630 # Avaliar o desempenho da RNA (LBP)
631 print("LBP")
632 print(classification_report(y_test_lbp, y_pred_rna_lbp, zero_division
    =1))
633 conf_matrix_rna_lbp = confusion_matrix(y_test_lbp, y_pred_rna_lbp)
634 specificity_rna_lbp = calculate_specificity(conf_matrix_rna_lbp)
635 print(f"Especificidade: {specificity_rna_lbp}")
636
637 # Avaliar o desempenho da RNA (LBP)
638 print("\n\nVGG16")
639 print(classification_report(y_test_vgg16, y_pred_rna_vgg16,
    zero_division=1))
640 conf_matrix_rna_vgg16 = confusion_matrix(y_test_vgg16,
    y_pred_rna_vgg16)
641 specificity_rna_vgg16 = calculate_specificity(conf_matrix_rna_vgg16)
642 print(f"Especificidade: {specificity_rna_vgg16}")
643
644 #RNA
645
646 #LBP
647 #           precision recall f1-score support
648
649 #         0         0.41    0.45    0.43         101
650 #         1         0.33    0.42    0.37          90
651 #         2         0.28    0.46    0.35          90
652 #         3         1.00    0.01    0.02          90
653
654 #   accuracy                   0.34         371
655 #  macro avg    0.51    0.33    0.29         371
656 #weighted avg    0.50    0.34    0.30         371
657
658 #Especificidade: 0.5357142857142857
659
660 #VGG16
661 #           precision recall f1-score support
662
663 #         0         0.92    0.95    0.94         101
664 #         1         0.81    0.58    0.68          90
665 #         2         0.69    0.73    0.71          90
666 #         3         0.84    1.00    0.91          90
667
668 #   accuracy                   0.82         371
669 #  macro avg    0.82    0.82    0.81         371
670 #weighted avg    0.82    0.82    0.81         371
671
672 #Especificidade: 0.9504950495049505
673

```

```

674 # Resultados: O modelo que deu melhores resultados foi o **RNA**,
    considerando a métrica de **sensibilidade** (recall).
675
676 ### RNA - Sensibilidade
677 #* CNN VGG16: **0.82**
678 #* LBP: **0.34**

```

```

1  import tensorflow as tf
2  from tensorflow.keras.preprocessing.image import ImageDataGenerator
3  from tensorflow.keras.applications import VGG16, ResNet50
4  from tensorflow.keras.layers import Dense, Flatten,
    GlobalAveragePooling2D
5  from tensorflow.keras.models import Model
6  from tensorflow.keras.optimizers import Adam
7  import numpy as np
8  import matplotlib.pyplot as plt
9  from sklearn.metrics import confusion_matrix, classification_report
10 import seaborn as sns
11 import os
12
13 img_height, img_width = 224, 224
14 batch_size = 32
15 num_classes = 4 # Classes: 0, 1, 2, 3
16
17 train_datagen = ImageDataGenerator(
18     preprocessing_function=tf.keras.applications.vgg16.
        preprocess_input # Pré-processamento específico
19 )
20
21 val_datagen = ImageDataGenerator(
22     preprocessing_function=tf.keras.applications.vgg16.
        preprocess_input
23 )
24
25 train_generator = train_datagen.flow_from_directory(
26     train_dir,
27     target_size=(img_height, img_width),
28     batch_size=batch_size,
29     class_mode='categorical'
30 )
31
32 validation_generator = val_datagen.flow_from_directory(
33     val_dir,
34     target_size=(img_height, img_width),
35     batch_size=batch_size,
36     class_mode='categorical'
37 )
38
39 def create_vgg16_model():
40     # Carregar o modelo base VGG16 sem as camadas totalmente
        conectadas (top)

```

```

41     base_model_vgg = VGG16(weights='imagenet', include_top=False,
42                               input_shape=(img_height, img_width, 3))
43
44     # Congelar as camadas convolucionais
45     for layer in base_model_vgg.layers:
46         layer.trainable = False
47
48     # Adicionar novas camadas
49     x = base_model_vgg.output
50     x = Flatten()(x)
51     x = Dense(1024, activation='relu')(x)
52     predictions = Dense(num_classes, activation='softmax')(x)
53
54     # Criar o modelo final
55     model_vgg = Model(inputs=base_model_vgg.input, outputs=predictions
56                       )
57
58     return model_vgg
59
60 model_vgg = create_vgg16_model()
61
62 model_vgg.compile(optimizer=Adam(learning_rate=1e-4),
63                  loss='categorical_crossentropy',
64                  metrics=['accuracy'])
65
66 epochs = 5
67
68 history_vgg = model_vgg.fit(
69     train_generator,
70     steps_per_epoch=train_generator.samples // batch_size,
71     epochs=epochs,
72     validation_data=validation_generator,
73     validation_steps=validation_generator.samples // batch_size
74 )
75
76 def create_resnet50_model():
77     base_model_resnet = ResNet50(weights='imagenet', include_top=False,
78                                   input_shape=(img_height, img_width, 3))
79     for layer in base_model_resnet.layers:
80         layer.trainable = False
81
82     x = base_model_resnet.output
83     x = GlobalAveragePooling2D()(x)
84     x = Dense(1024, activation='relu')(x)
85     predictions = Dense(num_classes, activation='softmax')(x)
86     model_resnet = Model(inputs=base_model_resnet.input, outputs=
87                           predictions)
88     return model_resnet
89
90 model_resnet = create_resnet50_model()
91
92 model_resnet.compile(optimizer=Adam(learning_rate=1e-4),

```

```

88         loss='categorical_crossentropy',
89         metrics=['accuracy'])
90
91 history_resnet = model_resnet.fit(
92     train_generator,
93     steps_per_epoch=train_generator.samples // batch_size,
94     epochs=epochs,
95     validation_data=validation_generator,
96     validation_steps=validation_generator.samples // batch_size
97 )
98
99 train_datagen_aug = ImageDataGenerator(
100     preprocessing_function=tf.keras.applications.vgg16.
101         preprocess_input,
102     rotation_range=20,
103     width_shift_range=0.1,
104     height_shift_range=0.1,
105     horizontal_flip=True
106 )
107
108 train_generator_aug = train_datagen_aug.flow_from_directory(
109     train_dir,
110     target_size=(img_height, img_width),
111     batch_size=batch_size,
112     class_mode='categorical'
113 )
114
115 # Treinamento com Data Augmentation
116 model_vgg_aug = create_vgg16_model()
117 model_vgg_aug.compile(optimizer=Adam(learning_rate=1e-4),
118     loss='categorical_crossentropy',
119     metrics=['accuracy'])
120
121 history_vgg_aug = model_vgg_aug.fit(
122     train_generator_aug,
123     steps_per_epoch=train_generator_aug.samples // batch_size,
124     epochs=epochs,
125     validation_data=validation_generator,
126     validation_steps=validation_generator.samples // batch_size
127 )
128
129 # Treinamento com Data Augmentation
130 model_resnet_aug = create_vgg16_model()
131 model_resnet_aug.compile(optimizer=Adam(learning_rate=1e-4),
132     loss='categorical_crossentropy',
133     metrics=['accuracy'])
134
135 history_resnet_aug = model_resnet_aug.fit(
136     train_generator_aug,
137     steps_per_epoch=train_generator_aug.samples // batch_size,
138     epochs=epochs,

```

```

138     validation_data=validation_generator,
139     validation_steps=validation_generator.samples // batch_size
140 )
141
142 test_datagen = ImageDataGenerator(
143     preprocessing_function=tf.keras.applications.vgg16.
144         preprocess_input
145 )
146 test_generator = test_datagen.flow_from_directory(
147     test_dir,
148     target_size=(img_height, img_width),
149     batch_size=1, # Para avaliar imagem por imagem
150     class_mode='categorical',
151     shuffle=False # Importante para manter a ordem das imagens
152 )
153
154 # VGG sem Data Augmentation
155 test_generator.reset()
156 preds_vgg = model_vgg.predict(test_generator, steps=test_generator.
157     samples)
158
159 # VGG com Data Augmentation
160 test_generator.reset()
161 preds_vgg_aug = model_vgg_aug.predict(test_generator, steps=
162     test_generator.samples)
163
164 # Converter as predições para classes
165 y_pred_vgg = np.argmax(preds_vgg, axis=1)
166 y_pred_vgg_aug = np.argmax(preds_vgg_aug, axis=1)
167 y_pred_resnet = np.argmax(preds_resnet, axis=1)
168 y_pred_resnet_aug = np.argmax(preds_resnet_aug, axis=1)
169
170 # Obter as classes verdadeiras
171 y_true = test_generator.classes
172 class_labels = list(test_generator.class_indices.keys())
173
174 from sklearn.metrics import confusion_matrix
175
176 # VGG16 sem Data Augmentation
177 cm_vgg = confusion_matrix(y_true, y_pred_vgg)
178
179 # VGG16 com Data Augmentation
180 cm_vgg_aug = confusion_matrix(y_true, y_pred_vgg_aug)
181
182 # ResNet50 sem Data Augmentation
183 cm_resnet = confusion_matrix(y_true, y_pred_resnet)
184
185 # ResNet50 com Data Augmentation
186 cm_resnet_aug = confusion_matrix(y_true, y_pred_resnet_aug)

```

```

186 from sklearn.metrics import classification_report, confusion_matrix
187
188 def calculate_specificity(cm):
189     # cm é a matriz de confusão
190     FP = cm.sum(axis=0) - np.diag(cm)
191     TN = cm.sum() - (FP + cm.sum(axis=1) - np.diag(cm) + np.diag(cm))
192     specificity = TN / (TN + FP)
193     return specificity
194
195 # Avaliar o desempenho do modelo VGG16 sem Data Augmentation
196 print("VGG16 sem Data Augmentation\n")
197 print(classification_report(y_true, y_pred_vgg, target_names=
198     class_labels))
199 conf_matrix_vgg = confusion_matrix(y_true, y_pred_vgg)
200 specificity_vgg = calculate_specificity(conf_matrix_vgg)
201 print(f"Especificidade: {specificity_vgg}")
202
203 # Avaliar o desempenho do modelo VGG16 com Data Augmentation
204 print("\nVGG16 com Data Augmentation\n")
205 print(classification_report(y_true, y_pred_vgg_aug, target_names=
206     class_labels))
207 conf_matrix_vgg_aug = confusion_matrix(y_true, y_pred_vgg_aug)
208 specificity_vgg_aug = calculate_specificity(conf_matrix_vgg_aug)
209 print(f"Especificidade: {specificity_vgg_aug}")
210
211 # Avaliar o desempenho do modelo ResNet50 sem Data Augmentation
212 print("\nResNet50 sem Data Augmentation\n")
213 print(classification_report(y_true, y_pred_resnet, target_names=
214     class_labels))
215 conf_matrix_resnet = confusion_matrix(y_true, y_pred_resnet)
216 specificity_resnet = calculate_specificity(conf_matrix_resnet)
217 print(f"Especificidade: {specificity_resnet}")
218
219 # Avaliar o desempenho do modelo ResNet50 com Data Augmentation
220 print("\nResNet50 com Data Augmentation\n")
221 print(classification_report(y_true, y_pred_resnet_aug, target_names=
222     class_labels))
223 conf_matrix_resnet_aug = confusion_matrix(y_true, y_pred_resnet_aug)
224 specificity_resnet_aug = calculate_specificity(conf_matrix_resnet_aug)
225 print(f"Especificidade: {specificity_resnet_aug}")
226
227 #VGG16 sem Data Augmentation
228
229 #           precision recall f1-score support
230
231 #         0         0.77    0.90    0.83     101
232 #         1         0.52    0.18    0.26      90
233 #         2         0.63    0.91    0.75      90
234 #         3         0.97    0.99    0.98      90

```

232	#	accuracy		0.75	371	
233	#	macro avg	0.72	0.74	0.70	371
234	#	weighted avg	0.72	0.75	0.71	371
235						
236	#	Especificidade:	[0.9	0.94661922	0.82918149	0.98932384]
237						
238	#	VGG16 com Data Augmentation				
239						
240	#	precision recall f1-score support				
241						
242	#	0	0.81	0.99	0.89	101
243	#	1	0.81	0.29	0.43	90
244	#	2	0.65	0.84	0.73	90
245	#	3	0.92	1.00	0.96	90
246						
247	#	accuracy		0.79	371	
248	#	macro avg	0.80	0.78	0.75	371
249	#	weighted avg	0.80	0.79	0.76	371
250						
251	#	Especificidade:	[0.91111111	0.97864769	0.85409253	0.97153025]
252						
253	#	ResNet50 sem Data Augmentation				
254						
255	#	precision recall f1-score support				
256						
257	#	0	0.99	0.97	0.98	101
258	#	1	0.93	0.56	0.69	90
259	#	2	0.69	0.97	0.81	90
260	#	3	0.98	1.00	0.99	90
261						
262	#	accuracy		0.88	371	
263	#	macro avg	0.90	0.87	0.87	371
264	#	weighted avg	0.90	0.88	0.87	371
265						
266	#	Especificidade:	[0.9962963	0.98576512	0.86120996	0.99288256]
267						
268	#	ResNet50 com Data Augmentation				
269						
270	#	precision recall f1-score support				
271						
272	#	0	0.75	0.99	0.85	101
273	#	1	0.58	0.12	0.20	90
274	#	2	0.61	0.81	0.70	90
275	#	3	0.91	1.00	0.95	90
276						
277	#	accuracy		0.74	371	
278	#	macro avg	0.71	0.73	0.68	371
279	#	weighted avg	0.71	0.74	0.68	371
280						
281	#	Especificidade:	[0.87407407	0.97153025	0.83629893	0.96797153]
282						

```
283 # Resultados: O modelo que deu melhores resultados foi o **ResNet50  
    sem Data Augmentation**, considerando a métrica de **sensibilidade  
    ** (recall).  
284  
285  
286 ### ResNet50 sem Data Augmentation - Sensibilidade  
287 #*   **0.88**
```

APÊNDICE 11 – ASPECTOS FILOSÓFICOS E ÉTICOS DA IA

A - ENUNCIADO

Título do Trabalho: "Estudo de Caso: Implicações Éticas do Uso do ChatGPT"

Trabalho em Grupo: O trabalho deverá ser realizado em grupo de alunos de no máximo seis (06) integrantes.

Objetivo do Trabalho: Investigar as implicações éticas do uso do ChatGPT em diferentes contextos e propor soluções responsáveis para lidar com esses dilemas.

Parâmetros para elaboração do Trabalho:

1. Relevância Ética: O trabalho deve abordar questões éticas significativas relacionadas ao uso da inteligência artificial, especialmente no contexto do ChatGPT. Os alunos devem identificar dilemas éticos relevantes e explorar como esses dilemas afetam diferentes partes interessadas, como usuários, desenvolvedores e a sociedade em geral.

2. Análise Crítica: Os alunos devem realizar uma análise crítica das implicações éticas do uso do ChatGPT em estudos de caso específicos. Eles devem examinar como o algoritmo pode influenciar a disseminação de informações, a privacidade dos usuários e a tomada de decisões éticas. Além disso, devem considerar possíveis vieses algorítmicos, discriminação e questões de responsabilidade.

3. Soluções Responsáveis: Além de identificar os desafios éticos, os alunos devem propor soluções responsáveis e éticas para lidar com esses dilemas. Isso pode incluir sugestões para políticas, regulamentações ou práticas de design que promovam o uso responsável da inteligência artificial. Eles devem considerar como essas soluções podem equilibrar os interesses de diferentes partes interessadas e promover valores éticos fundamentais, como transparência, justiça e privacidade.

4. Colaboração e Discussão: O trabalho deve envolver discussões em grupo e colaboração entre os alunos. Eles devem compartilhar ideias, debater diferentes pontos de vista e chegar a conclusões informadas através do diálogo e da reflexão mútua. O estudo de caso do ChatGPT pode servir como um ponto de partida para essas discussões, incentivando os alunos a aplicar conceitos éticos e legais aprendidos ao analisar um caso concreto.

5. Limite de Palavras: O trabalho terá um limite de 6 a 10 páginas teria aproximadamente entre 1500 e 3000 palavras.

6. Estruturação Adequada: O trabalho siga uma estrutura adequada, incluindo introdução, desenvolvimento e conclusão. Cada seção deve ocupar uma parte proporcional do total de páginas, com a introdução e a conclusão ocupando menos espaço do que o desenvolvimento.

7. Controle de Informações: Evitar incluir informações desnecessárias que possam aumentar o comprimento do trabalho sem contribuir significativamente para o conteúdo. Concentre-se em informações relevantes, argumentos sólidos e evidências importantes para apoiar sua análise.

8. Síntese e Clareza: O trabalho deverá ser conciso e claro em sua escrita. Evite repetições desnecessárias e redundâncias. Sintetize suas ideias e argumentos de forma eficaz para transmitir suas mensagens de maneira sucinta.

9. Formatação Adequada: O trabalho deverá ser apresentado nas normas da ABNT de acordo com as diretrizes fornecidas, incluindo margens, espaçamento, tamanho da fonte e estilo de citação. Deve-se seguir o seguinte template de arquivo: <https://bibliotecas.ufpr.br/wp-content/uploads/2022/03/template-artigo-de-periodico.docx>

B - RESOLUÇÃO

Estudo de Caso: Implicações Éticas do Uso do ChatGPT

RESUMO

Este estudo de caso investiga as implicações éticas do uso do ChatGPT em diferentes contextos e propõe soluções responsáveis para lidar com esses dilemas. A análise aborda questões éticas significativas, incluindo a disseminação de informações, privacidade dos usuários, vieses algorítmicos, discriminação e

responsabilidade. Soluções responsáveis são propostas para equilibrar os interesses das partes interessadas e promover valores fundamentais como transparência, justiça e privacidade.

Palavras-chave: Inteligência Artificial, Ética, ChatGPT, Privacidade, Vieses.

ABSTRACT This case study investigates the ethical implications of using ChatGPT in different contexts and proposes responsible solutions to address these dilemmas. The analysis addresses significant ethical issues, including information dissemination, user privacy, algorithmic biases, discrimination, and responsibility. Responsible solutions are proposed to balance stakeholder interests and promote fundamental values such as transparency, fairness, and privacy.

Keywords: Artificial Intelligence, Ethics, ChatGPT, Privacy, Biases.

1. INTRODUÇÃO

De acordo com Aruda (2024), em novembro de 2022, houve uma certa euforia e também assombro quando a empresa OpenAI tornou público e acessível a plataforma ChatGPT. Segundo a empresa OpenAI (2022), o ChatGPT é um modelo de inteligência artificial (IA) que interage de maneira conversacional, ou seja, ele foi programado para que haja um diálogo entre o usuário e a máquina, permitindo que a plataforma responda a perguntas, admita erros, conteste premissas incorretas e rejeite solicitações inadequadas.

Com a definição dada pelos próprios criadores da ferramenta, é possível identificar inúmeras lacunas relacionadas a questões éticas provenientes do processo de geração de respostas por meio de IA. Nesse sentido, alguns dilemas abordam problemáticas éticas relacionadas a: (1) privacidade e segurança dos dados; (2) desinformação e uso indevido; (3) viés discriminatório; (4) impactos nas formas de trabalho; (5) responsabilidade e transparência; (6) interações humanas; e, por último, (7) robustez e precisão das informações geradas automaticamente. A discussão acerca de cada item acima gera um cenário multifacetado, com perspectivas provenientes de diversas áreas. No âmbito da jurisprudência, o contexto torna-se ainda mais complexo, exigindo uma compreensão mais aprofundada para que as questões éticas sejam cuidadosamente examinadas.

Neste trabalho abordaremos três questões éticas, sendo elas: (1) privacidade e segurança dos dados; (2) desinformação e uso indevido e (7) robustez e precisão das informações geradas de forma automática. Quando se fala em privacidade dos dados no uso do ChatGPT, é importante ressaltar que a funcionalidade da ferramenta requer um grande volume de dados para o treinamento do modelo da IA por trás dela. Além disso, é necessário que as requisições sejam feitas pelo usuário. Desta forma, fica evidente que os dados são elementos protagonistas e participam de todas as camadas desta aplicação.

No entanto, muitos debates têm sido levantados acerca da responsabilização pela posse dos dados, sejam eles sensíveis ou não, que a plataforma possui. Uma vez que os dados são fornecidos como entrada, eles se tornam parte do sistema. Neste sentido, há abertura de brechas relacionadas ao uso e compartilhamento de dados sensíveis de forma indevida, causando danos quase que irreversíveis aos afetados. Segundo Wu et al. (2024) embora existam mecanismos de proteção de privacidade no ChatGPT, como o bloqueio de acesso a dados pessoais de indivíduos, não é garantido que não ocorra nenhum vazamento.

Aliado a isso, o modelo de IA (LLM - *Large Language Model*, tradução para o português: Modelo de Linguagem de Grande Escala) usado para a geração de informações pelo ChatGPT não é simples de entender, o que torna o processo pouco transparente e suscetível a gerar informações pouco precisas sobre determinado assunto e até informações falsas. Essa falta de transparência pode manipular a opinião pública, causando desinformação e influenciando decisões de forma negativa. Para elucidar estas questões o trabalho discutirá sobre cada um dos itens de forma mais aprofundada trazendo alguns exemplos reais dos impactos gerados pelos dilemas éticos.

2. DESENVOLVIMENTO

A revisão de literatura foca nos estudos prévios relacionados ao uso ético da IA e os desafios associados à implementação destes sistemas. Neste sentido, esta seção tem como objetivo trazer as definições necessárias presentes na literatura e os debates acerca dos impactos gerados por cada um dos dilemas éticos latentes no assunto.

2.1 Privacidade e segurança dos dados

Para Westin (1967): A privacidade na era digital é um direito fundamental que deve ser protegido com rigor. A coleta massiva de dados pessoais sem o consentimento adequado dos usuários representa uma ameaça significativa à autonomia e à dignidade individual.

A privacidade dos usuários é uma questão em evidência neste momento em que a IA está sendo inserida na sociedade e as pessoas começam a conhecer e utilizar. A partir do momento em que interagimos com o ChatGPT, ele armazena todas as informações em sua base de dados para treinamento e isso implica que esses dados se tornarão públicos, a partir do momento em que se faz uma pergunta e a ferramenta entender que aqueles dados inseridos por outra pessoa são os melhores para uma resposta a qualquer usuário. Então, surge a grande preocupação quanto a privacidade dos dados, que dados serão utilizados para a análise, que dados serão apresentados nas respostas, existe algum contato humano com esta base de dados? Não basta somente informar ao usuário de que ele deve evitar inserir dados sensíveis, e sim, o ChatGPT deve utilizar critérios éticos que aperfeiçoem a ferramenta para que ela analise esses dados e gere uma nova conclusão e resposta, considerando os critérios da ética e segurança e dando a eles o máximo de importância.

No estudo apresentado por Khowaja et. al (2024) são discutidas as preocupações acerca da privacidade dos dados usados para o treinamento de LLMs, em especial do ChatGPT. Os autores mencionam que os dados utilizados para o treinamento do ChatGPT são sistematicamente extraídos de mensagens, websites, artigos, livros e informações pessoais sem a devida autorização. Além disso, de acordo com a política de privacidade do ChatGPT (OpenAI, 2024b), são coletadas informações como dados de interação do usuário com o site, configurações e tipo de navegador, e endereço IP, juntamente com o tipo de conteúdo com o qual os usuários interagem no ChatGPT. Eles também coletam informações sobre atividades de navegação em diferentes sites e ao longo de um determinado período de tempo. A política de privacidade também afirma que "Além disso, de tempos em tempos, podemos analisar o comportamento geral e as características dos usuários de nossos serviços e compartilhar informações agregadas, como estatísticas gerais de usuários, com terceiros, publicar essas informações agregadas ou torná-las geralmente disponíveis. Podemos coletar informações agregadas através dos Serviços, através de cookies e por outros meios descritos nesta política de privacidade" e ainda que "em certas circunstâncias, podemos fornecer suas informações pessoais a terceiros sem aviso prévio, a menos que exigido por lei".

2.2 Desinformação e uso indevido

Segundo Wardle e Derakhshan (2017): A desinformação, especialmente em plataformas digitais, pode ter consequências graves para a sociedade. Ela pode manipular opiniões, influenciar decisões eleitorais e até mesmo causar pânico em situações de crise.

Atualmente, a internet é a principal fonte de informação. Desta forma, o desafio não está apenas em obter um conteúdo relevante, mas também em filtrar as informações incorretas da grande quantidade disponível. Antes do lançamento do ChatGPT, as pessoas utilizavam outros métodos para confirmar a veracidade das informações, avaliando o nível de conhecimento do criador do conteúdo pelo rigor da linguagem, correção do formato e a fonte do texto, sendo esses indicadores de confiabilidade. No entanto, ao utilizar o ChatGPT, é possível observar que esses indicadores são excelentes nesses aspectos, dando a falsa sensação de confiabilidade. Os usuários podem cair na armadilha de confiar cegamente no conteúdo gerado pelo ChatGPT depois de executar testes simples. Essa confiança cega pode levar a julgamentos errados devido a hábitos arraigados. Em áreas críticas, como documentos de pesquisa médica, informações básicas de experimentos em larga escala e conteúdo de políticas, fazer referência a informações errôneas do ChatGPT e tirar conclusões incorretas pode ter consequências imprevisíveis. Esses riscos não decorrem de fraude intencional, mas sim dos erros do ChatGPT, apesar de sua aparência inicialmente confiável (Wu et al., 2024).

Outro aspecto muito importante relacionado ao uso do ChatGPT é a geração de notícias falsas, que têm implicações éticas de bastante relevância social e política e já são um dos principais problemas que vêm sendo enfrentados no Brasil. Estes Modelos de Linguagem Grande (LLMs) são muito populares e podem ser facilmente acessados para a geração de informações falsas ou enganosas que podem facilmente serem disseminadas nas redes sociais e plataforma digitais, podendo ter com consequências graves, como

por exemplo a manipulação de eleições e a incitação à violências (Khowaja et. al, 2024). Até que ponto as informações geradas pelo ChatGPT são verdadeiras? Esse é um grande desafio, pelo motivo de que em muitos momentos as informações necessitam ser validadas, muitas vezes depende de decisões que ainda precisam ser formuladas e votadas pela sociedade, categorias e governos. O fato de que as informações geradas pelo ChatGPT são originadas de uma base de dados e esta base foi construída pelo próprio ChatGPT, leva a conclusão de que esta base de dados deve ser constantemente validada, se possível por um órgão neutro externo que controle a legitimidade das informações e as distribua de forma a guardar a ética e os direitos de todas as partes envolvidas. Atualmente existe a versão gratuita e a versão paga do ChatGPT, e a principal diferença entre elas é a atualização da base de dados e funcionalidades. Já foi constatado que em determinados questionamentos ao ChatGPT, a ferramenta informou erroneamente e isso é muito sério, pois, a maior fatia de utilização é pela opção gratuita, que é a com a base de dados desatualizada e que fornece grandes vulnerabilidades. Uma regulamentação que defina a forma que esses dados são validados é de fundamental importância para garantir a segurança e confiabilidade das informações.

2.3 Robustez e precisão das informações geradas

Para O'Neil (2016): A automação na tomada de decisões promete eficiência e objetividade, mas também apresenta riscos consideráveis. Algoritmos podem perpetuar vieses existentes nos dados de treinamento, levando a decisões injustas ou discriminatórias.

Outro desafio é a tomada de decisões automatizadas, que dá poder para a ferramenta decidir o que fazer com base nas regras que foram criadas por seus programadores. Um grande exemplo são as IAs de recrutamento e seleção, que fazem a triagem com base em critérios definidos e podem muitas vezes serem injustas. Com relação ao ChatGPT, eu posso desde solicitar a criação de um código em determinada linguagem que diante de um perfil pré-selecionado busque a opção ideal em sua base, posso pedir para que ele gere uma imagem como por exemplo pedir uma imagem de uma pessoa bonita, e ele irá criar uma imagem de uma pessoa branca, loira, de olhos claros, ou seja, seguindo critérios já definidos pelos programadores ou seguindo uma tendência das imagens buscadas na internet. Deixar que a IA faça a escolha sozinha oferece riscos, pois ela foi configurada para buscar a melhor solução em sua base de dados e ela pode estar repleta de vieses.

3. CONCLUSÃO

Através de discussões em grupo e colaboração entre os membros da equipe, chegamos a algumas propostas para lidar com as implicações éticas do uso do ChatGPT. Nossa análise destacou a importância de criar regulamentações claras para garantir que o uso da IA respeite a privacidade dos usuários e evite a disseminação de informações falsas. Essas regulamentações devem definir padrões rigorosos para a coleta, armazenamento e uso dos dados, assegurando a proteção das informações pessoais.

Além disso, identificamos a necessidade de promover a transparência nos processos de decisão do ChatGPT. Isso pode ser alcançado através da implementação de mecanismos de auditoria e monitoramento contínuo, permitindo verificar e corrigir vieses algorítmicos, reduzindo o risco de discriminação.

Durante nossas discussões, também enfatizamos a importância de práticas de design ético na criação e atualização dos modelos de IA. Envolver equipes multidisciplinares e incluir feedback de usuários e especialistas externos pode contribuir para a criação de um sistema mais robusto e justo.

Finalmente, destacamos a necessidade de educar os usuários sobre os potenciais riscos e limitações do ChatGPT. Programas de conscientização e treinamentos específicos podem capacitar os usuários a utilizarem a IA de forma crítica e responsável, reconhecendo a importância da verificação das informações fornecidas pela ferramenta.

Em resumo, através da colaboração entre os membros da equipe e discussões aprofundadas, conseguimos desenvolver soluções práticas para lidar com as implicações éticas do uso do ChatGPT. As soluções propostas incluem regulamentações claras, transparência nos processos de decisão, práticas de design ético e educação do usuário, promovendo um uso mais seguro e justo dessa tecnologia.

4. REFERÊNCIAS

Aruda, E. P. Inteligência Artificial Generativa No Contexto Da Transformação do Trabalho Docente. 2024. Disponível em: <<https://doi.org/10.1590/0102-469848078>>.

- Deng, J.; Lin, Y. The benefits and challenges of ChatGPT: An overview. *Frontiers in Computing and Intelligent Systems*, v. 2, n. 2, p. 81-83, 2022.
- Khowaja, S. A.; Khuwaja, P.; Dev, K.; Wang, W.; Nkenyereye, L. Chatgpt needs spade (sustainability, privacy, digital divide, and ethics) evaluation: A review. *Cognitive Computation*, p. 1-23, 2024.
- O'Neil, C. *Armas de Destruição Matemática: Como o Big Data Aumenta a Desigualdade e Ameaça a Democracia*. Nova York: Crown Publishing Group, 2016.
- OpenAI [internet]. Introducing ChatGPT. Disponível em: <<https://openai.com/index/chatgpt/>>. Acessado em 01 de Julho de 2024a.
- OpenAi [internet]. Privacy policy. Disponível em: <<https://openai.com/policies/privacy-policy/>>. Acesso em 8 de Julho de 2024b.
- Wardle, C.; Derakhshan, H. *Desordem da Informação: Rumo a uma Estrutura Interdisciplinar para Pesquisa e Formulação de Políticas*. Conselho da Europa, 2017.
- Westin, A. F. *Privacidade e Liberdade*. Nova York: Atheneum, 1967.
- Wu, X.; Ran, D.; Jianbing, N. Unveiling security, privacy, and ethical concerns of ChatGPT. *Journal of Information and Intelligence*, v. 2, n. 2, p. 102-115, 2024.

APÊNDICE 12 – GESTÃO DE PROJETOS DE IA

A - ENUNCIADO

1 Objetivo

Individualmente, ler e resumir – seguindo o template fornecido – um dos artigos abaixo:

AHMAD, L.; ABDELRAZEK, M.; ARORA, C.; BANO, M; GRUNDY, J. Requirements practices and gaps when engineering human-centered Artificial Intelligence systems. *Applied Soft Computing*. 143. 2023. DOI <https://doi.org/10.1016/j.asoc.2023.110421>

NAZIR, R.; BUCAIONI, A.; PELLICCIONE, P.; Architecting ML-enabled systems: Challenges, best practices, and design decisions. *The Journal of Systems & Software*. 207. 2024. DOI <https://doi.org/10.1016/j.jss.2023.111860>

SERBAN, A.; BLOM, K.; HOOS, H.; VISSER, J. Software engineering practices for machine learning – Adoption, effects, and team assessment. *The Journal of Systems & Software*. 209. 2024. DOI <https://doi.org/10.1016/j.jss.2023.111907>

STEIDL, M.; FELDERER, M.; RAMLER, R. The pipeline for continuous development of artificial intelligence models – Current state of research and practice. *The Journal of Systems & Software*. 199. 2023. DOI <https://doi.org/10.1016/j.jss.2023.111615>

XIN, D.; WU, E. Y.; LEE, D. J.; SALEHI, N.; PARAMESWARAN, A. Whither AutoML? Understanding the Role of Automation in Machine Learning Workflows. In *CHI Conference on Human Factors in Computing Systems (CHI'21)*, Maio 8-13, 2021, Yokohama, Japão. DOI <https://doi.org/10.1145/3411764.3445306>

2 Orientações adicionais

Escolha o artigo que for mais interessante para você. Utilize tradutores e o Chat GPT para entender o conteúdo dos artigos – caso precise, mas escreva o resumo em língua portuguesa e nas suas palavras. Não esqueça de preencher, no trabalho, os campos relativos ao seu nome e ao artigo escolhido. No template, você deverá responder às seguintes questões:

- Qual o objetivo do estudo descrito pelo artigo?
- Qual o problema/oportunidade/situação que levou a necessidade de realização deste estudo?
- Qual a metodologia que os autores usaram para obter e analisar as informações do estudo?
- Quais os principais resultados obtidos pelo estudo?

Responda cada questão utilizando o espaço fornecido no template, sem alteração do tamanho da fonte (Times New Roman, 10), nem alteração do espaçamento entre linhas (1.0). Não altere as questões do template. Utilize o editor de textos de sua preferência para preencher as respostas, mas entregue o trabalho em PDF. **B - RESOLUÇÃO**

Artigo: *Requirements practices and gaps when engineering human-centered Artificial Intelligence systems*

De acordo com os autores, o principal objetivo do estudo apresentado é identificar lacunas nas práticas atuais de engenharia de requisitos (RE) para IA (RE4IA). Os autores também apresentam duas perguntas de pesquisa a serem respondidas ao longo do estudo. A primeira está relacionada ao mapeamento das diretrizes existentes na pesquisa e na indústria para o desenvolvimento de IA centrada no usuário (CNU), bem como aos aspectos que ainda estão ausentes. A segunda pergunta trata da identificação de lacunas entre a pesquisa sobre RE4IA e as práticas aplicadas na indústria.

Qual o objetivo do estudo descrito pelo artigo?

A engenharia de software que utiliza IA é uma área emergente que apresenta inúmeros desafios. Grandes empresas utilizam diretrizes para auxiliar as equipes no desenvolvimento de sistemas de IA centrados no usuário. Entretanto, as práticas adotadas por profissionais para o desenvolvimento de sistemas em RE ainda são pouco relatadas. Os autores destacam que o foco principal da IA nos últimos anos tem sido a construção de sistemas tecnicamente sólidos. No entanto, ao priorizar os aspectos técnicos e negligenciar os aspectos humanos no desenvolvimento de sistemas de IA, podem surgir danos físicos e mentais. Há uma lacuna associada às abordagens centradas no usuário (CNU), o que pode resultar em sistemas de IA enviesados. Os autores também mencionam exemplos de vieses cometidos por sistemas de IA, relacionados a fatores como raça, gênero, cor, religião, estado de saúde e pessoas com deficiência. Recentemente, observa-se uma tendência em explorar aspectos centrados no ser humano no desenvolvimento de softwares com IA. Essa tendência é evidenciada por uma SLR, que revela que a maioria dos softwares ainda carece de uma abordagem centrada no ser humano em relação à modelagem de requisitos. Os aspectos abordados na SLR indicam que as pesquisas existentes em RE para IA têm se concentrado principalmente em temas como ética, confiança e explicabilidade, embora as avaliações empíricas sejam limitadas ou inexistentes. Além disso, nenhum dos estudos investigou os aspectos centrados no ser humano adotados pela indústria ou como eles devem ser abordados na RE.

Qual a metodologia que os autores usaram para obter e analisar as informações do estudo?

Ao longo da leitura, percebe-se que os autores realizaram uma pesquisa sistemática da literatura (SLR) sobre as diretrizes do Google, Apple, Microsoft e ML Canvas, com o objetivo de compreender como os sistemas dessas empresas abordam a interação entre o usuário e a IA. Os autores descrevem brevemente o funcionamento de cada uma dessas diretrizes e, em seguida, conduzem uma SLR em trabalhos que utilizam o RE4IA no período de 2010 a 2021. Esses trabalhos são categorizados em dois grupos: CNU e técnicos. Após esse mapeamento, os autores dividem os resultados em categorias e os analisam com base nos seis capítulos das diretrizes do Google, que abrangem as seguintes áreas: necessidade do usuário, necessidade do modelo, necessidade dos dados, feedback e controle, explicabilidade e confiança, erros e falhas. Após o mapeamento, os autores realizaram uma pesquisa com especialistas no desenvolvimento de sistemas de IA para compreender quais aspectos das diretrizes identificadas deveriam ser considerados na RE. Nesse contexto, elaboraram um questionário para investigar os desafios e limitações enfrentados pelos profissionais ao desenvolver sistemas de IA. O questionário também incluía perguntas sobre quais diretrizes em CNU, identificadas na SLR, esses especialistas consideravam ao construir um sistema de IA.

Quais os principais resultados obtidos pelo estudo?

Durante o mapeamento das pesquisas sobre diretrizes realizadas tanto na indústria quanto na academia para a RE, os autores identificaram algumas diferenças importantes. Enquanto a indústria se preocupa principalmente com os tipos e a diversidade dos dados, a academia foca na quantidade e na qualidade dos dados. Na pesquisa com especialistas da área de IA, os autores obtiveram resultados que destacam as áreas de atuação dos participantes, sendo a maioria composta por cientistas de dados e pesquisadores em IA. Outro aspecto interessante identificado foi a presença de especialistas atuando amplamente nos setores de educação e governo. Os autores também relataram os resultados obtidos em relação à forma como os requisitos são construídos pelos especialistas. A análise mostrou que a maior preocupação desses profissionais está em atender aos requisitos funcionais dos sistemas de IA. Sobre as ferramentas utilizadas para levantamento de requisitos, destacaram-se PowerPoint, Word, Email, Excel, Confluence, Jira, entre outras. Além disso, os autores abordaram as linguagens utilizadas pelos especialistas durante o processo de modelagem de requisitos, com o Caso de Uso sendo o formato mais frequente. Um ponto relevante levantado pelos especialistas foi a identificação de desafios recorrentes durante o processo de RE. O principal desafio apontado por eles durante a RE refere-se à “lacuna de viabilidade”, ou seja, as limitações dos modelos de IA em relação ao que podem ou não realizar. Os autores indicam que as diretrizes centradas no ser humano devem ser amplamente consideradas durante RE e que os especialistas também indicaram essa questão em suas respostas para o desenvolvimento de RE4IA.

APÊNDICE 13 – FRAMEWORKS DE INTELIGÊNCIA ARTIFICIAL

A - ENUNCIADO

1 Classificação (RNA)

Implementar o exemplo de Classificação usando a base de dados Fashion MNIST e a arquitetura RNA vista na aula **FRA - Aula 10 - 2.4 Resolução de exercício de RNA - Classificação**. Além disso, fazer uma breve explicação dos seguintes resultados:

- Gráficos de perda e de acurácia;
- Imagem gerada na seção “**Mostrar algumas classificações erradas**”, apresentada na aula prática.

Informações:

- **Base de dados:** Fashion MNIST Dataset
- **Descrição:** Um dataset de imagens de roupas, onde o objetivo é classificar o tipo de vestuário. É semelhante ao famoso dataset MNIST, mas com peças de vestuário em vez de dígitos.
- **Tamanho:** 70.000 amostras, 784 features (28x28 pixels).
- **Importação do dataset:** Copiar código abaixo.

```
1 data = tf.keras.datasets.fashion_mnist
2 (x_train, y_train), (x_test, y_test) = fashion_mnist.load_data()
```

2 Regressão (RNA)

Implementar o exemplo de Classificação usando a base de dados Wine Dataset e a arquitetura RNA vista na aula **FRA - Aula 12 - 2.5 Resolução de exercício de RNA - Regressão**. Além disso, fazer uma breve explicação dos seguintes resultados:

- Gráficos de avaliação do modelo (loss);
- Métricas de avaliação do modelo (pelo menos uma entre MAE, MSE, R^2).

Informações:

- **Base de dados:** Wine Quality
- **Descrição:** O objetivo deste dataset prever a qualidade dos vinhos com base em suas características químicas. A variável target (y) neste exemplo será o score de qualidade do vinho, que varia de 0 (pior qualidade) a 10 (melhor qualidade)
- **Tamanho:** 1599 amostras, 12 features.
- **Importação:** Copiar código abaixo.

```

1
2 url = "https://archive.ics.uci.edu/ml/machine-learning-databases/wine
   -quality/winequality-red.csv"
3 data = pd.read_csv(url, delimiter=';')
4
5 Dica 1. Para facilitar o trabalho, renomeie o nome das colunas para
   português, dessa forma:
6
7 data.columns = [
8     'acidez_fixa',      # fixed acidity
9     'acidez_volatil',   # volatile acidity
10    'acido_citrico',     # citric acid
11    'acucar_residual',   # residual sugar
12    'cloretos',          # chlorides
13    'dioxido_de_enxofre_livre', # free sulfur dioxide
14    'dioxido_de_enxofre_total', # total sulfur dioxide
15    'densidade',         # density
16    'pH',                # pH
17    'sulfatos',          # sulphates
18    'alcool',            # alcohol
19    'score_qualidade_vinho' # quality
20 ]
21
22 Dica 2. Separe os dados (x e y) de tal forma que a última coluna (í
   ndice -1), chamada score_qualidade_vinho, seja a variável target (
   y)

```

3 Sistemas de Recomendação

Implementar o exemplo de Sistemas de Recomendação usando a base de dados Base_livros.csv e a arquitetura vista na aula **FRA - Aula 22 - 4.3 Resolução do Exercício de Sistemas de Recomendação**. Além disso, fazer uma breve explicação dos seguintes resultados:

- Gráficos de avaliação do modelo (loss);
- Exemplo de recomendação de livro para determinado Usuário.

Informações:

- **Base de dados:** Base_livros.csv
- **Descrição:** Esse conjunto de dados contém informações sobre avaliações de livros (Notas), nomes de livros (Titulo), ISBN e identificação do usuário (ID_usuario)
- **Importação:** Base de dados disponível no Moodle (UFPR Virtual), chamada Base_livros (formato .csv).

4 Deepdream

Implementar o exemplo de implementação mínima de Deepdream usando uma imagem de um felino - retirada do site Wikipedia - e a arquitetura Deepdream vista na aula **FRA - Aula 23 - Prática Deepdream**. Além disso, fazer uma breve explicação dos seguintes resultados:

- Imagem onírica obtida por Main Loop;

- Imagem onírica obtida ao levar o modelo até uma oitava;
- Diferenças entre imagens oníricas obtidas com Main Loop e levando o modelo até a oitava.

Informações:

- **Base de dados:** https://commons.wikimedia.org/wiki/File:Felis_catus-cat_on_snow.jpg
- **Importação da imagem:** Copiar código abaixo.

```

1
2 url = "https://commons.wikimedia.org/wiki/Special:FilePath/
   Felis_catus-cat_on_snow.jpg"
3
4 Dica: Para exibir a imagem utilizando display (display.html) use o
   link https://commons.wikimedia.org/wiki/File:Felis_catus-
   cat_on_snow.jpg

```

B - RESOLUÇÃO

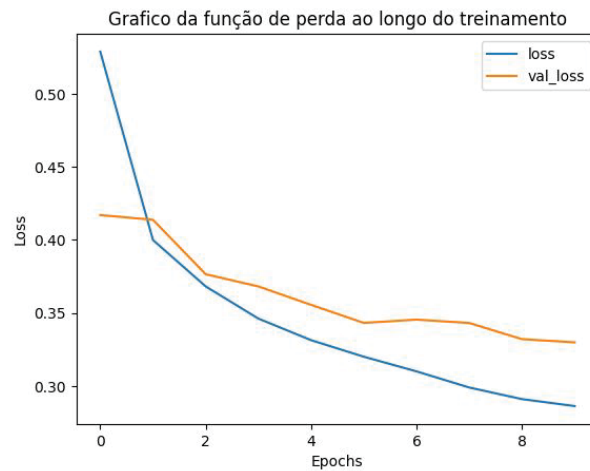
1 Classificação (RNA)

```

1 import tensorflow as tf
2 import numpy as np
3 import matplotlib.pyplot as plt
4 from mlxtend.plotting import plot_confusion_matrix
5 from sklearn.metrics import confusion_matrix
6
7 (X_train, Y_train), (X_test, Y_test) = tf.keras.datasets.
   fashion_mnist.load_data()
8
9 x_train, x_test = X_train/255.0, X_test/255.0
10
11 i = tf.keras.layers.Input(shape=(28,28))
12 x = tf.keras.layers.Flatten()(i)
13 x = tf.keras.layers.Dense(128, activation='relu')(x)
14 x = tf.keras.layers.Dropout(0.2)(x)
15 x = tf.keras.layers.Dense(10, activation='softmax')(x)
16 model = tf.keras.Model(i, x)
17
18 model.compile(optimizer='adam', loss='sparse_categorical_crossentropy',
   metrics=['accuracy'])
19
20 r = model.fit(x_train, Y_train, validation_data=(x_test, Y_test),
   epochs=10)
21
22 plt.plot(r.history['loss'], label='loss')
23 plt.plot(r.history['val_loss'], label='val_loss')
24 plt.xlabel('Epochs')
25 plt.ylabel('Loss')
26 plt.title('Gráfico da função de perda ao longo do treinamento')
27 plt.legend()

```

Figura 13.1 – EXERCÍCIO 1: GRÁFICO DE PERDA

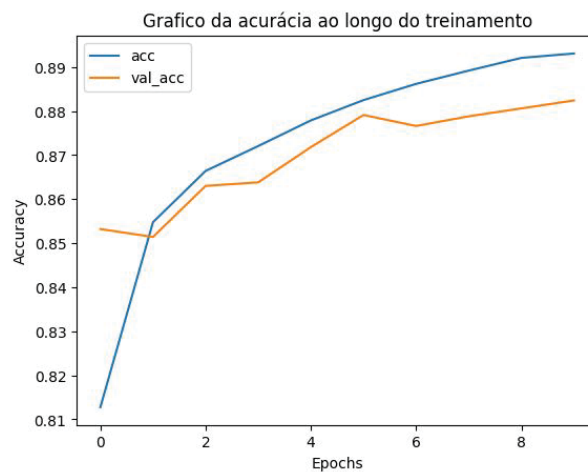


```

1 plt.plot(r.history['accuracy'], label='acc')
2 plt.plot(r.history['val_accuracy'], label='val_acc')
3 plt.xlabel('Epochs')
4 plt.ylabel('Accuracy')
5 plt.title('Gráfico da acurácia ao longo do treinamento')
6 plt.legend()

```

Figura 13.2 – EXERCÍCIO 1: GRÁFICO DE ACURÁCIA

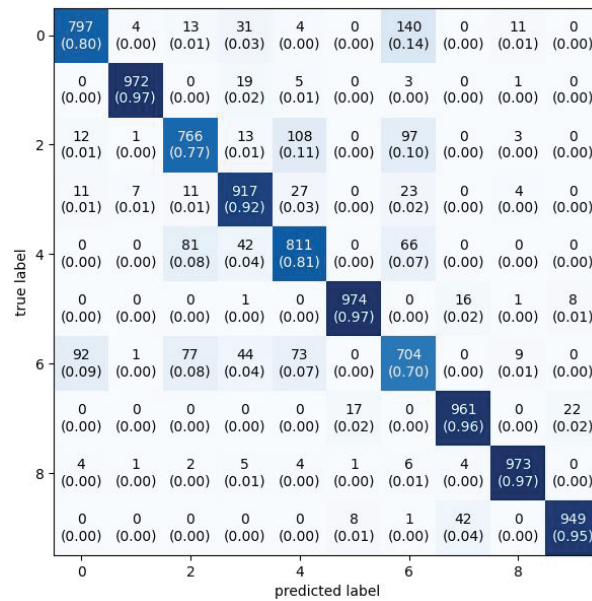


```

1 y_pred = model.predict(x_test).argmax(axis=1)
2 cm = confusion_matrix(Y_test, y_pred)
3 plot_confusion_matrix(conf_mat=cm, figsize=(7,7), show_normed=True)

```

Figura 13.3 – EXERCÍCIO 1: MATRIZ DE CONFUSÃO

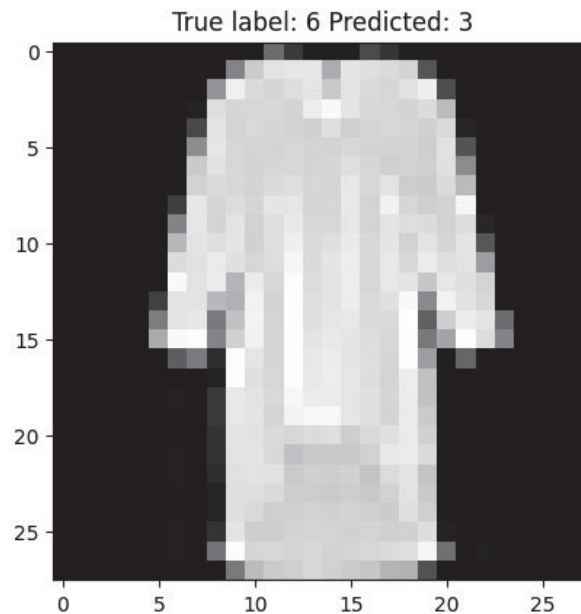


```

1 misclassified = np.where(y_pred != Y_test)[0]
2 i = np.random.choice(misclassified)
3 plt.imshow(x_test[i], cmap='gray')
4 plt.title("True label: %s Predicted: %s" % (Y_test[i], y_pred[i]));

```

Figura 13.4 – EXERCÍCIO 1: APRESENTANDO CLASSIFICAÇÃO INCORRETA



2 Regressão (RNA)

```

1 import tensorflow as tf
2 import numpy as np
3 import pandas as pd
4 import matplotlib.pyplot as plt

```

```

5 from tensorflow.python.keras import backend
6 from sklearn.model_selection import train_test_split
7 from sklearn.metrics import r2_score, mean_squared_error
8 from math import sqrt
9
10 url = "https://archive.ics.uci.edu/ml/machine-learning-databases/wine
    -quality/winequality-red.csv"
11
12 data = pd.read_csv(url, delimiter=';')
13
14 data.columns = [
15 'acidez_fixa', # fixed acidity
16 'acidez_volatil', # volatile acidity
17 'acido_citrico', # citric acid
18 'acucar_residual', # residual sugar
19 'cloretos', # chlorides
20 'dioxido_de_enxofre_livre', # free sulfur dioxide
21 'dioxido_de_enxofre_total', # total sulfur dioxide
22 'densidade', # density
23 'pH', # pH
24 'sulfatos', # sulphates
25 'alcool', # alcohol
26 'score_qualidade_vinho' # quality
27 ]
28
29 data = data.values
30
31 # Dica 2. Separe os dados (x e y) de tal forma que a última coluna (í
    ndice -1), chamada score_qualidade_vinho
32
33 X = data[:, 0:-1].astype(float)
34 Y = data[:, -1].astype(float)
35
36 X_train, X_test, Y_train, Y_test = train_test_split(X, Y, test_size
    =0.25)
37
38 i = tf.keras.layers.Input(shape=(11,))
39 x = tf.keras.layers.Dense(100, activation='relu')(i)
40 x = tf.keras.layers.Dense(1)(x)
41 model = tf.keras.Model(i, x)
42
43 # Criação de funções para as métricas r2 e RMSE serem inseridas no
    modelo
44 def rmse(y_true, y_pred):
45     return backend.sqrt(backend.mean(backend.square(y_true- y_pred),
        axis=-1))
46
47 def r2(y_true, y_pred):
48     media = backend.mean(y_true)
49     num = backend.sum(backend.square(y_true - y_pred))
50     den = backend.sum(backend.square(y_true - media))

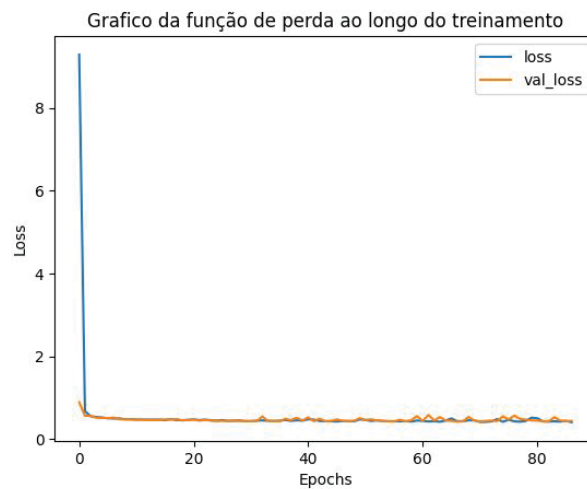
```

```

51     return (1.0 - num/den)
52
53 # Compilação
54 optimizer = tf.keras.optimizers.Adam(learning_rate=0.001)
55 model.compile(optimizer=optimizer, loss='mse', metrics=[r2, rmse])
56
57 # Early Stop para Epochs
58 early_stop = tf.keras.callbacks.EarlyStopping(monitor='val_loss',
59         patience=20, restore_best_weights=True)
60
61 r = model.fit(X_train, Y_train, validation_data=(X_test, Y_test),
62         epochs=1500, callbacks=[early_stop])
63
64 plt.plot(r.history['loss'], label='loss')
65 plt.plot(r.history['val_loss'], label='val_loss')
66 plt.xlabel('Epochs')
67 plt.ylabel('Loss')
68 plt.title('Gráfico da função de perda ao longo do treinamento')
69 plt.legend()

```

Figura 13.5 – EXERCÍCIO 2: GRÁFICO DE PERDA

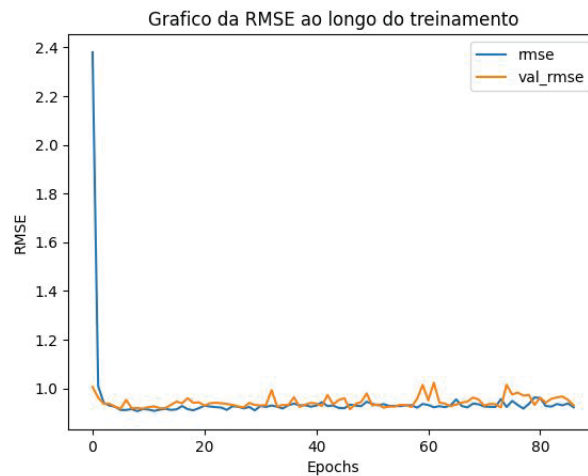


```

1 plt.plot(r.history['rmse'], label='rmse')
2 plt.plot(r.history['val_rmse'], label='val_rmse')
3 plt.xlabel('Epochs')
4 plt.ylabel('RMSE')
5 plt.title('Gráfico da RMSE ao longo do treinamento')
6 plt.legend()

```

Figura 13.6 – EXERCÍCIO 2: GRÁFICO DE RMSE

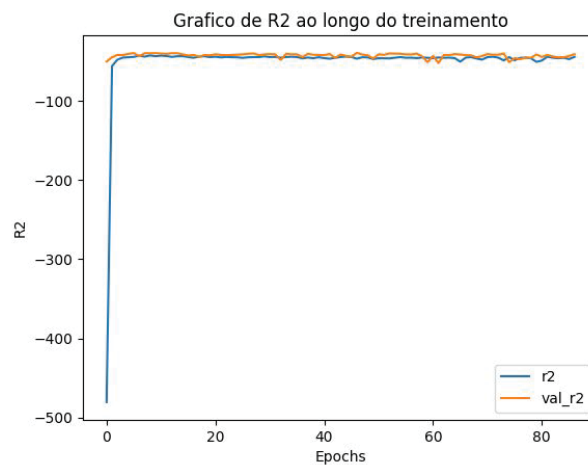


```

1 plt.plot(r.history['r2'], label='r2')
2 plt.plot(r.history['val_r2'], label='val_r2')
3 plt.xlabel('Epochs')
4 plt.ylabel('R2')
5 plt.title('Gráfico de R2 ao longo do treinamento')
6 plt.legend()

```

Figura 13.7 – EXERCÍCIO 2: GRÁFICO DE R2



```

1 y_pred = model.predict(X_test).flatten()
2 mse = mean_squared_error(Y_test, y_pred)
3 rmse = sqrt(mse)
4 r2 = r2_score(Y_test, y_pred)
5
6 print("mse:", mse)
7 print("rmse:", rmse)
8 print("r2:", r2)

```

mse: 0.4207298387142231
rmse: 0.6486369082269549

r2: 0.38145588854025325

3 Sistemas de Recomendação

```

1 import numpy as np
2 import pandas as pd
3 import tensorflow as tf
4 import matplotlib.pyplot as plt
5 from tensorflow.keras.models import Model
6 from tensorflow.keras.layers import Input, Embedding, Dense,
    Embedding, Flatten, Concatenate
7 from tensorflow.keras.optimizers import SGD, Adam
8 from sklearn.utils import shuffle
9
10 df = pd.read_csv('Base_livros(in).csv', sep=',', on_bad_lines='skip')
11 df['Notas;;;;;'] = df['Notas;;;;;'].str.replace(r'^0-9', '',
    regex=True).astype(float)
12 df.rename(columns={'Notas;;;;;': 'Notas'}, inplace=True)
13 df.dropna(inplace=True)
14
15 categorical_columns = ['ID_usuario', 'Titulo', 'Autor', 'Editora']
16 for cat_column in categorical_columns:
17     df[cat_column] = pd.Categorical(df[cat_column])
18     df['new_' + cat_column] = df[cat_column].cat.codes
19
20 N = len(set(df['ID_usuario']))
21 M = len(set(df['new_Titulo']))
22
23 # dimensão do embedding
24 K = 10
25
26 # usuário
27 u = Input(shape=(1,))
28 u_emb = Embedding(N, K)(u) # saída : num_samples, 1, K
29 u_emb = Flatten()(u_emb) # saída : num_samples, K
30 # filme
31 m = Input(shape=(1,))
32 m_emb = Embedding(M, K)(m) # saída : num_samples, 1, K
33 m_emb = Flatten()(m_emb) # saída : num_samples, K
34 x = Concatenate()([u_emb, m_emb])
35 x = Dense(1024, activation="relu")(x)
36 x = Dense(1)(x)
37 model = Model(inputs=[u, m], outputs=x)
38
39 model.compile(loss="mse", optimizer=SGD(learning_rate=0.08, momentum
    =0.9))
40
41 user_ids, movie_ids, ratings = shuffle(df['new_ID_usuario'], df['
    new_Titulo'], df['Notas'])
42 Ntrain = int(0.8 * len(ratings)) # separar os dados 80% x 20%
43 train_user = user_ids[:Ntrain]
44 train_movie = movie_ids[:Ntrain]

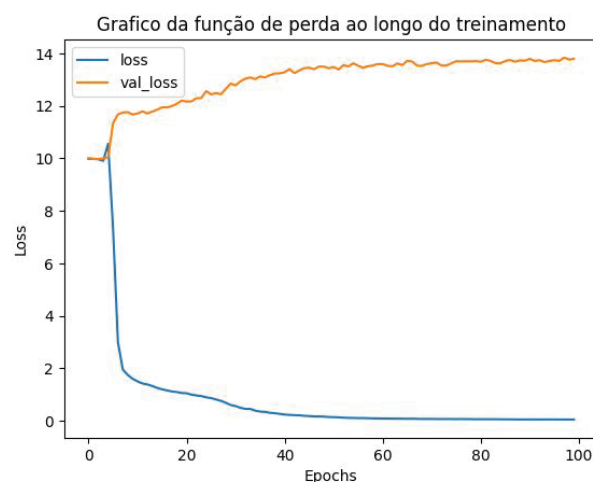
```

```

45 train_ratings = ratings[:Ntrain]
46 test_user = user_ids[Ntrain:]
47 test_movie = movie_ids[Ntrain:]
48 test_ratings = ratings[Ntrain:]
49
50 # centralizar as notas
51 avg_rating = train_ratings.mean()
52 train_ratings = train_ratings - avg_rating
53 test_ratings = test_ratings - avg_rating
54
55 epochs = 100
56 r = model.fit(
57 x=[train_user, train_movie],
58 y=train_ratings,
59 epochs=epochs,
60 batch_size=1024,
61 verbose=2, # não imprime o progresso
62 validation_data=(test_user, test_movie, test_ratings)
63 )
64
65 plt.plot(r.history['loss'], label='loss')
66 plt.plot(r.history['val_loss'], label='val_loss')
67 plt.xlabel('Epochs')
68 plt.ylabel('Loss')
69 plt.title('Grafico da função de perda ao longo do treinamento')
70 plt.legend()

```

Figura 13.8 – EXERCÍCIO 3: GRÁFICO DE PERDA



```

1 input_usuario = np.repeat(a=8263, repeats=M)
2 film = np.array(list(set(movie_ids)))
3 preds = model.predict([input_usuario, film])
4
5 # descentraliza as predições
6 rat = preds.flatten() + avg_rating
7 # índice da maior nota

```

```

8 idx = np.argmax(rat)
9 print("Recomendação: Filme - ", film[idx], " / ", rat[idx] , "*")

```

Recomendação: Filme - 66366 / 19.101202 *

4 Deepdream

```

1 import tensorflow as tf
2 import numpy as np
3 import matplotlib as mpl
4 import IPython.display as display
5 import PIL.Image
6
7 url = "https://commons.wikimedia.org/wiki/Special:FilePath/
8     Felis_catus-cat_on_snow.jpg"
9
10 # Download da imagem e gravação em array Numpy
11 def download(url, max_dim=None):
12     name = url.split('/')[-1]
13     image_path = tf.keras.utils.get_file(name, origin=url)
14     img = PIL.Image.open(image_path)
15     if max_dim:
16         img.thumbnail((max_dim, max_dim))
17     return np.array(img)
18
19 # Normalização da imagem
20 def deprocess(img):
21     img = 255*(img + 1.0)/2.0
22     return tf.cast(img, tf.uint8)
23
24 # Mostra a imagem
25 def show(img):
26     display.display(PIL.Image.fromarray(np.array(img)))
27
28 # Redução do tamanho da imagem para facilitar o trabalho da RNN
29 original_img = download(url, max_dim=500)
30 show(original_img)

```

Figura 13.9 – EXERCÍCIO 3: IMAGEM BAIXADA DA URL FORNECIDA



```

1 base_model = tf.keras.applications.InceptionV3(include_top=False,
2     weights='imagenet')
3 # Maximizando as ativações das camadas
4 names = ['mixed3', 'mixed5']
5 layers = [base_model.get_layer(name).output for name in names]
6
7 # Criação do modelo
8 dream_model = tf.keras.Model(inputs=base_model.input, outputs=layers)
9
10 def calc_loss(img, model):
11     # Passe a imagem pelo modelo para recuperar as ativações.
12     # Converte a imagem em um batch de tamanho 1.
13     img_batch = tf.expand_dims(img, axis=0)
14     layer_activations = model(img_batch)
15
16     if len(layer_activations) == 1:
17         layer_activations = [layer_activations]
18
19     losses = []
20     for act in layer_activations:
21         loss = tf.math.reduce_mean(act)
22         losses.append(loss)
23
24     return tf.reduce_sum(losses)
25
26 class DeepDream(tf.Module):
27     def __init__(self, model):
28         self.model = model
29         @tf.function(
30             input_signature=(
31                 tf.TensorSpec(shape=[None, None, 3], dtype=tf.float32),
32                 tf.TensorSpec(shape=[], dtype=tf.int32),
33                 tf.TensorSpec(shape=[], dtype=tf.float32),)
34             )
35         def __call__(self, img, steps, step_size):
36             print("Tracing")
37             loss = tf.constant(0.0)
38
39             for n in tf.range(steps):
40                 with tf.GradientTape() as tape:
41                     # Gradientes relativos a img
42                     tape.watch(img)
43                     loss = calc_loss(img, self.model)
44
45                     # Calculo do gradiente da perda em relação aos pixels da
46                     # imagem de entrada.
47                     gradients = tape.gradient(loss, img)
48
49                     # Normalizacao dos gradintes
50                     gradients /= tf.math.reduce_std(gradients) + 1e-8

```

```

50
51     # Na subida gradiente, a "perda" é maximizada.
52     # Você pode atualizar a imagem adicionando diretamente os
      gradientes (porque eles têm o mesmo f
53     img = img + gradients*step_size
54     img = tf.clip_by_value(img, -1, 1)
55
56     return loss, img
57
58 deepdream = DeepDream(dream_model)
59
60 def run_deep_dream_simple(img, steps=100, step_size=0.01):
61     img = tf.keras.applications.inception_v3.preprocess_input(img)
62     img = tf.convert_to_tensor(img)
63     step_size = tf.convert_to_tensor(step_size)
64     steps_remaining = steps
65     step = 0
66     while steps_remaining:
67         if steps_remaining>100:
68             run_steps = tf.constant(100)
69         else:
70             run_steps = tf.constant(steps_remaining)
71
72         steps_remaining -= run_steps
73         step += run_steps
74
75         loss, img = deepdream(img, run_steps, tf.constant(step_size))
76
77         display.clear_output(wait=True)
78         show(deprocess(img))
79         print ("Step {}, loss {}".format(step, loss))
80
81     result = deprocess(img)
82     display.clear_output(wait=True)
83     show(result)
84
85     return result
86
87 dream_img = run_deep_dream_simple(img=original_img, steps=100,
      step_size=0.01)

```

Figura 13.10 – EXERCÍCIO 3: RESULTADO OBTIDO NO PROCESSO DE DEEP DREAM



```

1 import time
2 start = time.time()
3
4 OCTAVE_SCALE = 1.30
5 img = tf.constant(np.array(original_img))
6 base_shape = tf.shape(img)[: -1]
7 float_base_shape = tf.cast(base_shape, tf.float32)
8 for n in range(-2, 3):
9     new_shape = tf.cast(float_base_shape*(OCTAVE_SCALE**n), tf.int32)
10    img = tf.image.resize(img, new_shape).numpy()
11    img = run_deep_dream_simple(img=img, steps=50, step_size=0.01)
12
13 display.clear_output(wait=True)
14 img = tf.image.resize(img, base_shape)
15 img = tf.image.convert_image_dtype(img/255.0, dtype=tf.uint8)
16 show(img)
17 end = time.time()
18 end-start

```

Figura 13.11 – EXERCÍCIO 3: RESULTADO OBTIDO NO PROCESSO DE DEEP DREAM À OITAVA



APÊNDICE 14 – VISUALIZAÇÃO DE DADOS E STORYTELLING

A - ENUNCIADO

Escolha um conjunto de dados brutos (ou uma visualização de dados que você acredite que possa ser melhorada) e faça uma visualização desses dados (de acordo com os dados escolhidos e com a ferramenta de sua escolha) Desenvolva uma narrativa/storytelling para essa visualização de dados considerando os conceitos e informações que foram discutidas nesta disciplina. Não esqueça de deixar claro para seu possível público alvo qual **o objetivo dessa visualização de dados, o que esses dados significam, quais possíveis ações podem ser feitas com base neles.**

Entregue em um PDF:

- **O conjunto de dados brutos (ou uma visualização de dados** que você acredite que possa ser melhorada);
- Explicação do **contexto e o público-alvo** da visualização de dados e do storytelling que será desenvolvido;
- **A visualização desses dados** (de acordo com os dados escolhidos e com a ferramenta de sua escolha) **explicando a escolha do tipo de visualização e da ferramenta usada;** (50 pontos)

B - RESOLUÇÃO

1) O conjunto de dados brutos (ou uma visualização de dados que você acredite que possa ser melhorada);

Devido ao grande volume de linhas que compõem o arquivo com os dados brutos do conjunto de dados utilizado nesta atividade, esses dados foram posicionados ao final deste documento.

2) Explicação do contexto e o público-alvo da visualização de dados e do storytelling que será desenvolvido;

O conjunto de dados escolhido para esta atividade consiste em um relatório contendo a listagem de todos os filmes, séries e programas disponíveis na Netflix. Os dados foram obtidos por meio da plataforma Kaggle, disponível neste link. Segundo a descrição do conjunto de dados, a Netflix é uma das plataformas de streaming de mídia mais populares, com mais de 8 mil filmes e programas de TV disponíveis. Em 2024, a plataforma contava com mais de 282 milhões de assinantes em todo o mundo.

O público-alvo para a visualização destes dados pode abranger diferentes grupos, dentre eles podemos destacar:

- Pesquisadores que possuem interesse em compreender as tendências de consumo, análise de representatividade, impacto na indústria do entretenimento;
- Profissionais do marketing que possuem interesse em compreender as preferências do público (popularidade) para criar campanhas segmentadas;
- Criadores de conteúdo que possuem interesse em compartilhar assuntos baseados na tendência;
- Críticos do cinema que possuem interesse em analisar padrões de produção e prever o sucesso ou fracasso do gênero;
- Empresas que queiram concorrer com a Netflix no sentido de compreender como é realizado o engajamento da plataforma com seu público;
- Investidores que possuem interesse em analisar as tendências do setor de streaming e visam enxergar oportunidades de investimentos;

- Jornalistas e especialistas em entretenimento que possuem interesse em relatar sobre a plataforma e realizar comparações com outras soluções do mercado. Sendo estes, também, fortemente relacionados com os criadores de conteúdos citados acima.

As colunas do conjunto de dados são:

- show_id: trata-se de um identificador do item dentro do conjunto;
- type: trata-se de uma especificação entre as categorias a qual a mídia pertence, podendo ser movie (filme) ou TV show (programa de TV);
- title: o nome do título dentro da Netflix;
- director: diretor do título;
- country: país qual o título foi produzido;
- date_added: corresponde a data a qual o título foi adicionado a plataforma;
- release_year: corresponde ao ano que o título foi lançado;
- rating: classificação do conteúdo do título;
- duration: duração do título em minutos;
- listed_in: categorias ou gêneros em que o título se enquadra;
- description: uma breve descrição sobre o título.

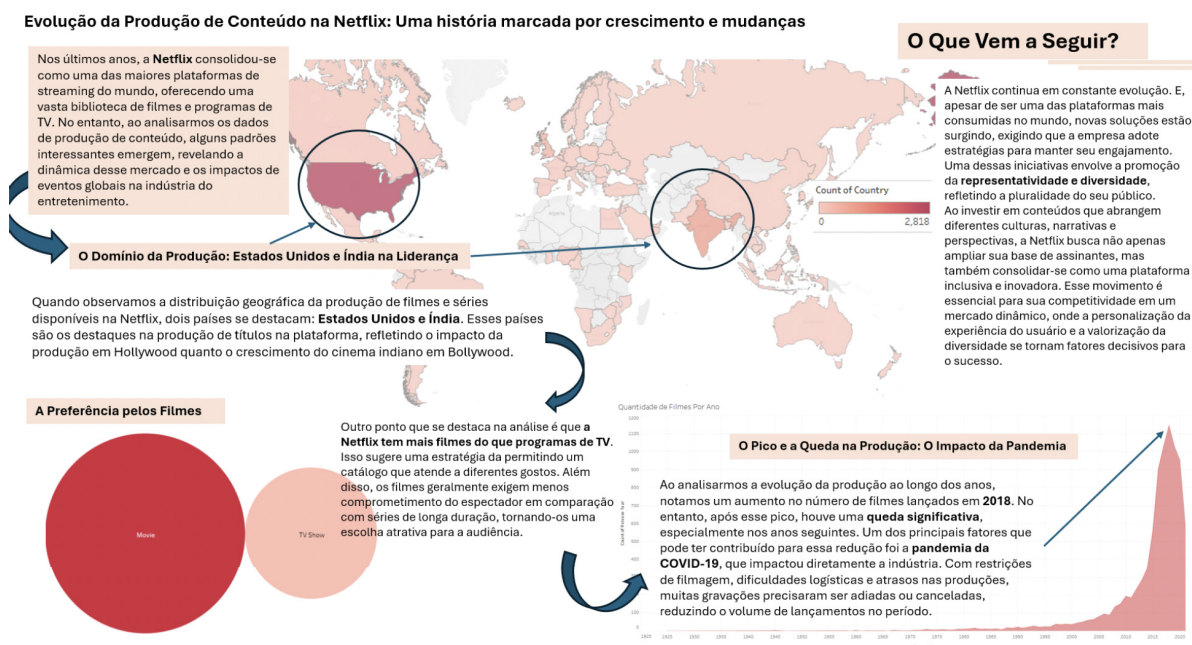
3) A visualização desses dados (de acordo com os dados escolhidos e com a ferramenta de sua escolha) explicando a escolha do tipo de visualização e da ferramenta usada; (50 pontos);

A ferramenta Public Tableau foi selecionada para gerar as visualizações deste trabalho. Ela foi escolhida devido à sua facilidade de uso e à variedade de opções e sugestões automáticas que fornece. Neste trabalho foram escolhidas algumas visualizações para descrever o perfil dos títulos presentes no Netflix. As visualizações escolhidas foram:

- Quantidade de títulos por país, indicando a geolocalização. A prevalência é representada por uma escala de cor vermelha (escolhida por remeter à identidade visual da plataforma), onde tons mais intensos indicam um maior número de filmes produzidos. Nesta visualização, é possível perceber que os Estados Unidos e a Índia são os dois países com o maior número de títulos na Netflix.
- Distribuição por tipos de mídia dentro da plataforma, distinguindo por tamanho o volume de itens que pertencem as categorias filme e programas de TV;
- Quantidade de filmes lançados por ano indicando o crescimento e a queda na produção podendo ser relacionado com acontecimentos no mundo.

4) A descrição da narrativa/storytelling dessa visualização de dados. (50 pontos).

Figura 14.1 – EVOLUÇÃO DA PRODUÇÃO DO CONTEÚDO NA NETFLIX



APÊNDICE 15 – TÓPICOS EM INTELIGÊNCIA ARTIFICIAL

A - ENUNCIADO

1) Algoritmo Genético

Problema do Caixeiro Viajante

A Solução poderá ser apresentada em: Python (preferencialmente), ou em R, ou em Matlab, ou em C ou em Java. Considere o seguinte problema de otimização (a escolha do número de 100 cidades foi feita simplesmente para tornar o problema intratável. A solução ótima para este problema não é conhecida). Suponha que um caixeiro deva partir de sua cidade, visitar clientes em outras 99 cidades diferentes, e então retornar à sua cidade. Dadas as coordenadas das 100 cidades, descubra o percurso de menor distância que passe uma única vez por todas as cidades e retorne à cidade de origem.

Para tornar a coisa mais interessante, as coordenadas das cidades deverão ser sorteadas (aleatórias), considere que cada cidade possui um par de coordenadas (x e y) em um espaço limitado de 100 por 100 pixels. O relatório deverá conter no mínimo a primeira melhor solução (obtida aleatoriamente na geração da população inicial) e a melhor solução obtida após um número mínimo de 1000 gerações. Gere as imagens em 2d dos pontos (cidades) e do caminho.

Sugestão:

- considere o cromossomo formado pelas cidades, onde a cidade de início (escolhida aleatoriamente) deverá estar na posição 0 e 100 e a ordem das cidades visitadas nas posições de 1 a 99 deverão ser definidas pelo algoritmo genético.
- A função de avaliação deverá minimizar a distância euclidiana entre as cidades (os pontos).
- Utilize no mínimo uma população com 100 indivíduos;
- Utilize no mínimo 1% de novos indivíduos obtidos pelo operador de mutação;
- Utilize no mínimo de 90% de novos indivíduos obtidos pelo método de cruzamento (crossover-ox);
- Preserve sempre a melhor solução de uma geração para outra.

Importante: A solução deverá implementar os operadores de “cruzamento” e “mutação”.

2) Compare a representação de dois modelos vetoriais

Pegue um texto relativamente pequeno, o objetivo será visualizar a representação vetorial, que poderá ser um vetor por palavra ou por sentença. Seja qual for a situação, considere a quantidade de palavras ou sentenças onde tenha no mínimo duas similares e no mínimo 6 textos, que deverão produzir no mínimo 6 vetores. Também limite o número máximo, para que a visualização fique clara e objetiva.

O trabalho consiste em pegar os fragmentos de texto e codificá-las na forma vetorial. Após obter os vetores, imprima-os em figuras (plot) que demonstrem a projeção desses vetores usando a PCA.

O PDF deverá conter o código-fonte e as imagens obtidas.

B - RESOLUÇÃO

1) Algoritmo Genético

```
1 import math
2 import random
3 import matplotlib.pyplot as plt
4
5 # Definindo parâmetros do problema
```

```

6 NUM_CIDADES = 100
7 TAMANHO_POPULACAO = 100
8 NUM_GERACOES = 5000
9
10 # Função para calcular a distancia total de uma rota (caminho fechado
    que passa
11 def calcular_distancia_total(rota, coordenadas):
12     """
13     Calcula a distância total percorrida seguindo a sequência de
        cidades em 'rota'.
14     A distância é a soma das distâncias euclidianas entre cada par
        consecutivo.
15     *Importante: A rota começa e termina na cidade de origem (rota[0])
        .
16     """
17     distancia = 0.0
18     # Soma as distâncias entre cada cidade e a próxima na rota
19     for i in range(len(rota) - 1):
20         cidade_atual = rota[i]
21         proxima_cidade = rota[i+1]
22         # Obtém as coordenadas das duas cidades
23         x1, y1 = coordenadas[cidade_atual]
24         x2, y2 = coordenadas[proxima_cidade]
25         # Calcula a distancia Euclidiana entre cidade atual e proxima
            cidade
26         distancia += math.hypot(x2-x1, y2-y1)
27     return distancia
28
29 # Função de seleção por torneio para escolher um individuo (rota)
    como pai
30 def selecionar_pai_por_torneio(populacao, distancias, k=3):
31     """
32     Seleciona e retorna um individuo da população usando seleção por
        torneio.
33     Aleatoriamente escolhe 'k' indivíduos e retorna o que tiver a
        menor distância.
34     """
35     melhor_indice = -1
36     melhor_distancia = float('inf')
37     # Realiza o torneio entre individuos aleatórios
38     for i in range(k):
39         indice = random.randrange(len(populacao)) # escolhe um indice
            aleatorio
40         # Verifica se este individuo é o melhor do torneio até agora
41         if distancias[indice] < melhor_distancia:
42             melhor_distancia = distancias[indice]
43             melhor_indice = indice
44         # Retorna uma cópia da rota vencedora (melhor distância)
45     return list(populacao[melhor_indice])
46
47 def crossover_OX(pai1, pai2):

```

```

48 | """
49 | Realiza o crossover OX (Order Crossover) entre dois cromossomos (
    | pais).
50 | Retorna um cromossomo filho válido, com cidade inicial fixa.
51 | """
52 | tamanho = len(pai1)
53 | # Inicializa o filho garantindo cidade inicial e final fixas
54 | filho = [None] * tamanho
55 | filho[0] = pai1[0]
56 | filho[-1] = pai1[-1]
57 |
58 | # Escolhe dois pontos de corte válidos e diferentes
59 | corte1, corte2 = sorted(random.sample(range(1, tamanho - 1), 2))
60 |
61 | # Copia segmento do pai1 diretamente para o filho
62 | segmento_pai1 = pai1[corte1:corte2 + 1]
63 | filho[corte1:corte2+1] = segmento_pai1
64 |
65 | # Guarda as cidades já inseridas no filho
66 | cidades_no_filho = set(segmento_pai1)
67 |
68 | # Preenche o restante das posições com cidades de pai2 (na ordem
    | correta)
69 | pos_filho = (corte2 + 1) % (tamanho - 1)
70 | pos_pai2 = (corte2 + 1) % (tamanho - 1)
71 |
72 | # Repete até o filho ser totalmente preenchido (sem None)
73 | while None in filho[1:-1]:
74 |     cidade = pai2[pos_pai2]
75 |     if cidade not in cidades_no_filho:
76 |         filho[pos_filho] = cidade
77 |         cidades_no_filho.add(cidade)
78 |         pos_filho = pos_filho + 1 if pos_filho + 1 < tamanho - 1
    | else 1
79 |         pos_pai2 = pos_pai2 + 1 if pos_pai2 + 1 < tamanho - 1 else 1
80 |
81 | return filho
82 |
83 | # Função de mutação que troca duas cidades de lugar na rota (mutação
    | por troca)
84 | def mutacao(rota):
85 |     """
86 |     Realiza uma mutação simples em uma rota, trocando de posição duas
    | cidades aleatórias.
87 |     A cidade de origem (posição 0 e última posição) não é alterada.
88 |     Retorna uma nova rota mutada.
89 |     """
90 |     # Cria uma cópia da rota para mutação (para não alterar a rota
    | original diretamente)
91 |     rota_mutada = list(rota)
92 |

```

```

93     # Seleciona aleatoriamente duas posições diferentes na parte
        intermediaria
94     i = random.randint(1, len(rota_mutada) - 2)
95     j = random.randint(1, len(rota_mutada) - 2)
96
97     while i == j:
98         j = random.randint(1, len(rota_mutada) - 2)
99
100    # Troca as cidades das posições i e j
101    rota_mutada[i], rota_mutada[j] = rota_mutada[j], rota_mutada[i]
102    return rota_mutada
103
104    # Gera coordenadas aleatórias para cada cidade dentro de um espaço
        100x100
105    coordenadas = [] # Lista para armazenar tuples (x, y) de coordenadas
        de cada cidade
106    for i in range(NUM_CIDADES):
107        x = random.random() * 100 # coordenada x aleatoria entre 0 e 100
108        y = random.random() * 100 # coordenada y aleatoria entre 0 e 100
109        coordenadas.append((x, y))
110    print("Coordenadas:", coordenadas)
111
112    # Gera a população inicial de rotas aleatórias
113    populacao = []
114    distancias = []
115    cidade_origem = 0 # definindo a cidade de indice 0 como cidade de
        origem
116
117    # Inicializa as variaveis para acompanhar a melhor solução encontrada
118    melhor_rota = None
119    melhor_distancia = float('inf')
120
121    for i in range(TAMANHO_POPULACAO):
122        # Cria uma ordem aleatória das cidades (exceto a cidade de origem)
123        outras_cidades = list(range(1, NUM_CIDADES))
124        random.shuffle(outras_cidades)
125        # Monta a rota completa: cidade de origem no início, seguida das
            outras cidades
126        rota = [cidade_origem] + outras_cidades + [cidade_origem]
127        # Adiciona a rota gerada à população
128        populacao.append(rota)
129        # Calcula a distância dessa rota e guarda
130        distancias.append(calcular_distancia_total(rota, coordenadas))
131
132    # Seleciona uma rota inicial aleatória (por exemplo, a primeira da
        população)
133    rota_inicial = populacao[0]
134    distancia_inicial = distancias[0]
135    print("Distancia inicial:", distancia_inicial)
136

```

```

137 # Identifica o melhor indivíduo da população inicial (com menor
    distancia)
138 for i, rota in enumerate(populacao):
139     d = distancias[i]
140     if d < melhor_distancia:
141         melhor_distancia = d
142         melhor_rota = rota
143
144 # Loop principal do Algoritmo Genético
145 for geracao in range(NUM_GERACOES):
146     nova_populacao = []
147     nova_distancias = []
148
149     # Elitismo: preserva a melhor rota da geração atual na próxima
150     nova_populacao.append(melhor_rota)
151     nova_distancias.append(melhor_distancia)
152
153     # Determina quantos novos indivíduos serão gerados por crossover e
        por mutação
154     num_por_crossover = int(TAMANHO_POPULACAO * 0.90)
155     num_por_mutacao = int(TAMANHO_POPULACAO * 0.09)
156     if num_por_mutacao < 1:
157         num_por_mutacao = 1
158
159     total_novos = 1 + num_por_crossover + num_por_mutacao # 1 elitismo
        + novos
160     if total_novos > TAMANHO_POPULACAO:
161         num_por_crossover -= (total_novos - TAMANHO_POPULACAO)
162     elif total_novos < TAMANHO_POPULACAO:
163         num_por_crossover += (TAMANHO_POPULACAO - total_novos)
164
165     # Gera novos indivíduos por crossover
166     for _ in range(num_por_crossover):
167         pai1 = selecionar_pai_por_torneio(populacao, distancias, k=3)
168         pai2 = selecionar_pai_por_torneio(populacao, distancias, k=3)
169         filho = crossover_OX(pai1, pai2)
170         nova_populacao.append(filho)
171         nova_distancias.append(calcular_distancia_total(filho,
            coordenadas))
172
173     # Gera novos indivíduos por mutação
174     for _ in range(num_por_mutacao):
175         base = selecionar_pai_por_torneio(populacao, distancias, k=3)
176         individuo_mutado = mutacao(base)
177         nova_populacao.append(individuo_mutado)
178         nova_distancias.append(calcular_distancia_total(
            individuo_mutado, coordenadas))
179
180     # Se por algum motivo excedemos o tamanho da população, cortamos o
        excedente
181     if len(nova_populacao) > TAMANHO_POPULACAO:

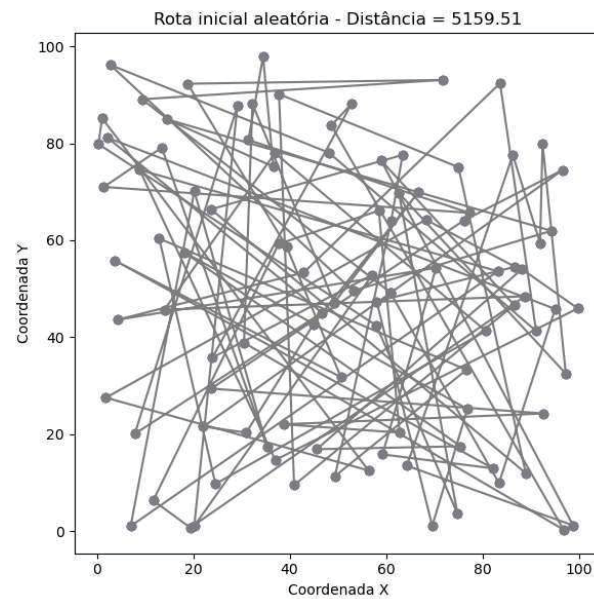
```

```

182     nova_populacao = nova_populacao[:TAMANHO_POPULACAO]
183     nova_distancias = nova_distancias[:TAMANHO_POPULACAO]
184
185     # Atualiza a população para a próxima geração
186     populacao = nova_populacao
187     distancias = nova_distancias
188
189     # Atualiza o melhor indivíduo
190     for i, d in enumerate(distancias):
191         if d < melhor_distancia:
192             melhor_distancia = d
193             melhor_rota = populacao[i]
194
195     # A cada 100 gerações, imprime a evolução
196     if (geracao + 1) % 100 == 0 or geracao == 0:
197         print(f"Geração {geracao + 1}: melhor distância encontrada {
198             melhor_distancia}")
199
200 # Imprime as distâncias da rota inicial e da melhor rota final
201 encontrada
202 print(f"\nDistancia total da rota inicial: {distancia_inicial:.2f}")
203 print(f"Distância total da melhor rota final: {melhor_distancia:.2f}\
204     n")
205
206 # Gera o gráfico da rota inicial aleatória
207 plt.figure(figsize=(6,6))
208 xs = [coordenadas[cidade][0] for cidade in rota_inicial]
209 ys = [coordenadas[cidade][1] for cidade in rota_inicial]
210 plt.plot(xs, ys, marker='o', color='gray')
211 plt.scatter(xs, ys, color="blue")
212 plt.title(f"Rota inicial aleatória - Distancia = {distancia_inicial
213     :.2f}")
214 plt.xlabel("Coordenada X")
215 plt.ylabel("Coordenada Y")
216 plt.tight_layout()

```

Figura 15.1 – EXERCÍCIO 1: ROTA INICIAL ALEATÓRIA, DISTÂNCIA = 5159,51

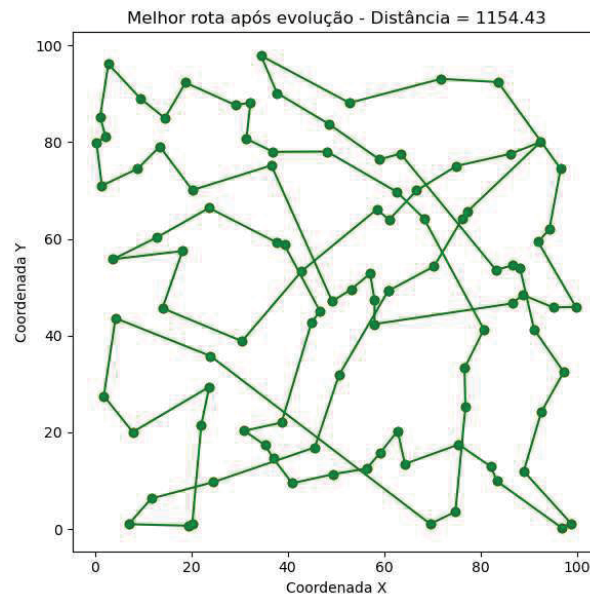


```

1 # Gera o gráfico da melhor rota final após a evolução
2 plt.figure(figsize=(6,6))
3 xs_best = [coordenadas[cidade][0] for cidade in melhor_rota]
4 ys_best = [coordenadas[cidade][1] for cidade in melhor_rota]
5 plt.plot(xs_best, ys_best, marker='o', color='green')
6 plt.scatter(xs_best, ys_best, color='red')
7 plt.title(f"Melhor rota após evolução - Distância = {melhor_distancia
8           :.2f}")
9 plt.xlabel("Coordenada X")
10 plt.ylabel("Coordenada Y")
11 plt.tight_layout()
12 # Mostra os gráficos gerados
13 plt.show()

```

Figura 15.2 – EXERCÍCIO 1: ROTA APÓS EVOLUÇÃO, DISTÂNCIA = 1154,43



2) Compare a representação de dois modelos vetoriais

```

1 # Comandos para manipulação de pacotes
2 !pip uninstall numpy -y
3 !pip install numpy==1.25.2
4 !pip uninstall gensim -y
5 !pip install --no-binary :all: gensim
6
7 import numpy as np
8 import re
9 from gensim.models import Word2Vec
10 from sklearn.decomposition import PCA
11 from scipy.spatial.distance import pdist, squareform
12 from scipy.cluster.hierarchy import dendrogram, linkage
13 import matplotlib.pyplot as plt
14
15 # Textos curtos de exemplo
16 textos = [
17     "O gato está dormindo no sofá.", # Texto 1
18     "O cachorro está dormindo no tapete.", # Texto 2 (similar to Texto
19         1)
20     "A tecnologia esta mudando nossas vidas.", # Texto 3
21     "Estamos vivendo na era da tecnologia.", # Texto 4 (similar to
22         Texto 3)
23     "O dia está ensolarado e claro.", # Texto 5
24     "O sol está brilhando no céu azul." # Texto 6 (similar to Texto 5)
25 ]
26
27 # Pre-processamento: impor pontuação e criar lista de tokens por
28     texto
29 sentencas_tokens = []
30 for txt in textos:
31     # Converter para minúsculas e remover pontuação

```

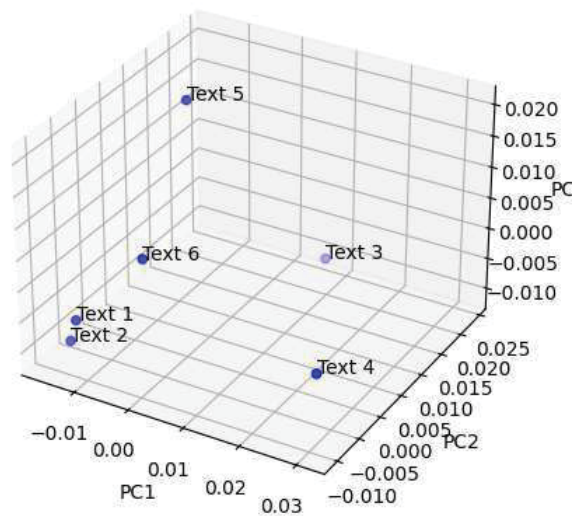
```

29     texto_limpo = re.sub(r'[\w\s]', '', txt.lower())
30     tokens = texto_limpo.split()
31     if tokens:
32         sentencas_tokens.append(tokens)
33
34 # Treinar modelo Word2Vec nos tokens
35 model = Word2Vec(sentencas_tokens, vector_size=50, window=5,
36                 min_count=1, workers=4)
37
38 # Obter vetor de cada sentença fazendo média dos vetores de palavras
39 vetores_texto = []
40 for tokens in sentencas_tokens:
41     vetores_palavras = [model.wv[word] for word in tokens if word in
42                        model.wv.key_to_index]
43     if len(vetores_palavras) > 0:
44         # Média dos vetores de palavra para representar a sentença
45         vetor_sentenca = np.mean(vetores_palavras, axis=0)
46     else:
47         # Caso sentença sem palavras conhecidas
48         vetor_sentenca = np.zeros(model.vector_size)
49     vetores_texto.append(vetor_sentenca)
50 vetores_texto = np.array(vetores_texto)
51
52 # Calcular vetor medio e centralizar os vetores de texto
53 vetor_medio = np.mean(vetores_texto, axis=0)
54 vetores_centralizados = vetores_texto - vetor_medio
55
56 # Aplicar PCA para reduzir dimensão dos vetores
57 pca_3d = PCA(n_components=3)
58 pca_2d = PCA(n_components=2)
59 vetores_3d = pca_3d.fit_transform(vetores_centralizados)
60 vetores_2d = pca_2d.fit_transform(vetores_centralizados)
61
62 # Visualização dos vetores em 3D e 2D
63 labels = [f"Text {i+1}" for i in range(len(textos))]
64
65 # Plot 3D dos vetores originais (após PCA 3D)
66 fig = plt.figure(figsize=(6,5))
67 ax = fig.add_subplot(111, projection='3d')
68 ax.scatter(vetores_3d[:,0], vetores_3d[:,1], vetores_3d[:,2], c='blue',
69           marker='o')
70 for i, label in enumerate(labels):
71     ax.text(vetores_3d[i,0], vetores_3d[i,1], vetores_3d[i,2], label)
72 ax.set_title('Vetores originais (projeção PCA 3D)')
73 ax.set_xlabel('PC1')
74 ax.set_ylabel('PC2')
75 ax.set_zlabel('PC3')
76 plt.show()

```

Figura 15.3 – EXERCÍCIO 2: PROJEÇÃO PCA 3D

Vetores originais (projeção PCA 3D)

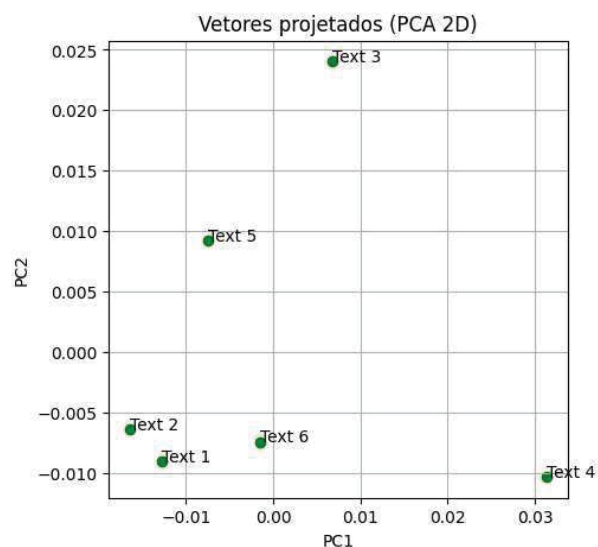


```

1 # Plot 2D dos vetores projetados (PCA 2D)
2 plt.figure(figsize=(5,5))
3 plt.scatter(vetores_2d[:,0], vetores_2d[:,1], c='green', marker='o')
4 for i, label in enumerate(labels):
5     plt.annotate(label, (vetores_2d[i,0], vetores_2d[i,1]))
6 plt.title('Vetores projetados (PCA 2D)')
7 plt.xlabel('PC1')
8 plt.ylabel('PC2')
9 plt.grid(True)
10 plt.show()

```

Figura 15.4 – EXERCÍCIO 2: PROJEÇÃO PCA 2D



```

1 # Calcular matrizes de distancia entre vetores

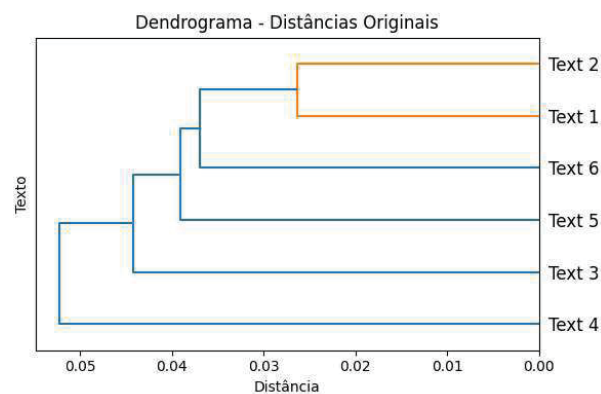
```

```

2 dist_original = squareform(pdist(vetores_centralizados, metric='
    euclidean'))
3 dist_projetado = squareform(pdist(vetores_2d, metric='euclidean'))
4
5 # Gerar dendrogramas de agrupamento hierárquico
6 # Dendrograma para distâncias originais
7 Z_orig = linkage(vetores_centralizados, method='ward', metric='
    euclidean')
8 plt.figure(figsize=(6,4))
9 dendrogram(Z_orig, labels=labels, orientation='left')
10 plt.title('Dendrograma - Distâncias Originais')
11 plt.xlabel('Distância')
12 plt.ylabel('Texto')
13 plt.tight_layout()
14 plt.show()

```

Figura 15.5 – EXERCÍCIO 2: DENDOGRAMA - DISTÂNCIAS ORIGINAIS



```

1 # Dendrograma para distâncias após projeção PCA 2D
2 Z_proj = linkage(vetores_2d, method='ward', metric='euclidean')
3 plt.figure(figsize=(6,4))
4 dendrogram(Z_proj, labels=labels, orientation='left')
5 plt.title('Dendrograma - Distâncias PCA 2D')
6 plt.xlabel('Distância')
7 plt.ylabel('Texto')
8 plt.tight_layout()
9 plt.show()

```

Figura 15.6 – EXERCÍCIO 2: DENDOGRAMA - DISTÂNCIAS PCA 2D

