

Universidade Federal de Uberlândia
Faculdade de Computação

XII FACOM TECHWEEK E XIX WORKSHOP DE TESES E DISSERTAÇÕES EM CIÊNCIA DA COMPUTAÇÃO

Anais

24 a 28 de novembro de 2025

ISSN: 2447-0406



Uberlândia
2025

XII FACOM TechWeek

Anais do XIX WTDCC 2025

24 a 28 de novembro de 2025

Comissão Organizadora

Adriano Mendonça Rocha (UFU)

Claudiney Ramos Tinoco (UFU)

Juliete Aparecida Ramos Costa (UFU/PPGCO)

Renan Gonçalves Cattelan (UFU)

Renata dos Santos Melo (UFU/PPGCO)

Thiago Pirola Ribeiro (UFU)

Organização, Execução, Promoção e Realização

Universidade Federal de Uberlândia (UFU)

Faculdade de Computação (FACOM/UFU)

FACOM Techweek e WTDCC

A TechWeek é um evento realizado pela Faculdade de Computação (FACOM) da Universidade Federal de Uberlândia (UFU), que busca promover atualização técnica e apresentar as últimas tendências tecnológicas, reunindo e proporcionando uma maior interação entre discentes, docentes, profissionais e empresas das áreas de Tecnologia da Informação. Trata-se de um evento tradicional, já em sua décima segunda edição, composto por palestras, minicursos, mesas-redondas, competições técnicas e apresentação de trabalhos e pesquisas científicas.

Desde 2016, o Workshop de Teses e Dissertações em Ciência da Computação (WTDCC) é realizado em conjunto com a TechWeek. Em sua décima nona edição, o WTDCC contou com a participação de discentes de graduação e de pós-graduação de universidades de Uberlândia e da região.

Dentre os objetivos do WTDCC destacam-se:

- Permitir a troca de experiência entre os pesquisadores;
- Promover a melhoria dos trabalhos produzidos, por meio de uma avaliação crítica e independente;
- Divulgar o Programa de Pós-Graduação em Computação da FACOM/UFU;
- Permitir aos docentes um acompanhamento das pesquisas em andamento;
- Apoiar os trabalhos de pesquisa dos discentes;
- Estimular a colaboração entre os docentes do PPGCO/UFU;
- Permitir a colaboração inter-grupos de pesquisa; e,
- Refletir sobre a atual situação do PPGCO na UFU e no cenário nacional.

Nesse sentido, o evento constitui um importante espaço de integração para a troca de experiências acadêmico-científicas, objetivando o desenvolvimento da ciência e da tecnologia nesse domínio do conhecimento.

A TechWeek e o WTDCC tiveram atividades desenvolvidas em formato presencial e online em 2025.

Prefácio

O XIX Workshop de Teses e Dissertações em Ciência da Computação (WTDCC) foi realizado juntamente com a XII FACOM TechWeek e teve como objetivo divulgar os projetos de pesquisa científica e tecnológica na área da computação. São projetos realizados por discentes da graduação e pós-graduação da Universidade Federal de Uberlândia (UFU) e da região. Além disso, teve como objetivo contribuir para a formação dos participantes, despertando o interesse pelas descobertas científicas e pela resolução de problemas complexos.

Na sua décima nona edição, o WTDCC contou com palestras temáticas, mesa-redonda, bem como com a submissão e apresentação de trabalhos. As palestras abrangeram tópicos relacionados às linhas de pesquisa desenvolvidas pelo Programa de Pós-Graduação em Ciência da Computação (PPGCO/UFU). A mesa-redonda reuniu os discentes internacionais do PPGCO/UFU para uma conversa sobre suas pesquisas e sua cultura. As submissões de trabalhos, em formato de resumo, contemplaram duas trilhas relacionadas à formação dos discentes: trabalhos de graduação (iniciação científica e trabalho de conclusão de curso) e de pós-graduação (especialização, mestrado e doutorado). Vários desses trabalhos foram apresentados durante o evento em sessões de exposição de pôsteres.

Além da publicação de 89 trabalhos, o WTDCC apresentou uma expressiva participação do público, interessado em conhecer mais sobre as pesquisas desenvolvidas em Uberlândia e região.

Nesse sentido, agradecemos ao público que prestigiou as palestras e as apresentações de trabalhos, às palestrantes convidadas, aos participantes da mesa redonda e a todos que submeteram seus trabalhos para o evento.

Por fim, agradecemos à comissão de avaliação dos trabalhos e também às empresas e instituições de Uberlândia e região que, de alguma forma, marcaram presença no evento.

Prof. Claudiney Ramos Tinoco
Juliete Aparecida Ramos Costa
Renata dos Santos Melo
Coordenadores do XIX WTDCC

Organização

Coordenação da TechWeek

Prof. Adriano Mendonça Rocha (UFU)
Prof. Renan Gonçalves Cattelan (UFU)
Prof. Thiago Pirola Ribeiro (UFU)

Coordenação do WTDCC

Prof. Claudiney Ramos Tinoco (UFU)
Juliete Aparecida Ramos Costa (UFU/PPGCO)
Renata dos Santos Melo (UFU/PPGCO)

Comissão de Avaliação dos Trabalhos do WTDCC

Abdelkader Tagmouni (UFU/PPGCO)
Profa. Alessandra Aparecida Paulino (UFU)
Prof. Alexsandro Santos Soares (UFU)
Prof. Caio Augusto Rodrigues dos Santos (UFU)
Prof. Claudiney Ramos Tinoco (UFU)
Daniel Ricardo Cunha Oliveira (UFU/PPGCO)
Prof. Daniel Stefany Duarte Caetano (UFU)
Prof. Diego Nunes Molinos (UFU)
Eduardo Cassiano da Silva (UFU/PPGCO)
Eduardo dos Santos Rocha (UFU/PPGCO)
Eneia Gazite (UFU/PPGCO)
Profa. Fabíola Souza Fernandes Pereira (UFU)
Giullia Rodrigues de Menezes (UFU/PPGCO)
Prof. Igor da Penha Natal (UFU)
Prof. Ivan Sendin (UFU)
Prof. José Gustavo de Souza Paiva (UFU)
Prof. Leonardo Muttoni (UFU)
Prof. Luiz Cláudio Theodoro (UFU)
Prof. Marcelo Zanchetta do Nascimento (UFU)
Prof. Maurício Cunha Escarpinati (UFU)
Prof. Paulo Henrique Ribeiro Gabriel (UFU)
Renata dos Santos Melo (UFU/PPGCO)
Prof. Rivalino Matias (UFU)
Profa. Roberta Barbosa Oliveira (UnB)
Prof. Roberto Júnior Silva Caetano (UFU)
Prof. Ronaldo Castro de Oliveira (UFU)
Saide Manuel Saide (UFU/PPGCO)
Salomão Bento Nilo Pena (UFU/PPGCO)
Prof. Shiguelo Nomura (UFU)
Prof. Stéphane Julia (UFU)
Prof. Thiago Fialho de Queiroz Lafetá (UFU)
Prof. Thiago Pirola Ribeiro (UFU)
Prof. Victor Sobreira (UFU)

Equipe de Trabalho

Adriano Mendonça Rocha
Alexandre Magno Silva Júnior
Ana Clara Alves Reis Silva
Ana Vitória Vaz Santos
Anna Clara Rodrigues Peres
Annelise Lima Carneiro
Arisa Nishigori Okada
Augusto Idalgo Costa
Bernardo Tibaldi Dias
Breno Melo Moreira
Bruno Ferrari Lacerra
Carla Azevedo Silva
César Reis Cardoso Júnior
Claudiney Ramos Tinoco
Daniel Stefany Duarte Caetano
Diego Nunes Molinos
Eduarda Lopes Santos Moura
Eduardo dos Santos Rocha
Eduardo Lordão Oliveira
Emílio da Silva Machado
Erick Lucius Felix
Evelyn Ariel Cardoso Barbosa
Felipe Roza Bonetti
Fernando Luiz de Paula Santil
Gabriel Antonio Martins Vieira
Gabriel Fernandes de Jesus
Guilherme Cabral de Menezes
Guilherme Epifânio da Silva
Gustavo Henrique Leal da Fonseca
Gustavo Luis De Siqueira Nascimento
Humberto Luiz Razente
Igor da Penha Natal
João Antônio Menezes Jordão
João Gabriel Santos Rodrigues
João Gabriel Tévez Dias Pereira
João Guilherme Araújo Viana
Juliete Aparecida Ramos Costa
Kauan Rodrigues Barros
Lavínia Barbosa Dantas de Souza
Leandro Henrique Silva Rabelo
Leonardo Cardoso de Moura
Luana Rodrigues Borges
Lucas Matos Rodrigues
Lucas Rodrigues Silva Flores

Luis Felipe Garcia de Souza Paim
Marco Antônio Caetano de Alcântara
Marcos Paulo Gomes Pires
Maria Eduarda Ferreira Teixeira
Maria Rita Vieira Souza
Marya Ysabella Colatino de Souza
Matheus Gualter Silva Resende
Matheus Marques Mendes Pereira
Nicolly Ribeiro Luz
Pablo Rodrigues Cardoso Araújo
Pablo Vinícius Carrijo
Pedro Guagneli Furlan
Pedro Miguel de Paula Silva
Rafael Alexandre Carneiro de Lima
Rafael Barcelos de Oliveira Reis
Rafael Pasquini
Rebecca da Silva Souza
Renan Cesar Vieira Mallagoli
Renan Gonçalves Cattelan
Renata dos Santos Melo
Sara Luzia de Melo
Tatiana Mayumi Tamura
Thiago Pirola Ribeiro
Victor Lourençato Brizante
Victor Sobreira
Vinicius Silva Tadeu
Vinicius Tavares Martins
Vitória Fernandes Costa Silva

Editoração dos Anais

Prof. Claudiney Ramos Tinoco (UFU)
Renata dos Santos Melo (UFU/PPGCO)

Sumário

1. Trilha: Trabalhos de Graduação	12
Aleatoriedade Quântica em Sistemas Fotônicos: Da Simulação Computacional à Criptografia Moderna	13
Algoritmo Genético Orientado à Desordem para a Otimização da Exploração de Ambientes por Enxames de Robôs	14
Análise Bibliométrica da Inteligência Artificial na Agricultura: Mapeamento de Tendências e Frentes de Pesquisa	15
Análise de Comentários de Vídeos do YouTube utilizando Redes Textuais	16
Análise de Diferentes Operações de Aumento de Dados em Modelos de Classificação de Imagens Histológicas	17
Análise de Resolvedores de DNS para Navegação Segura de Usuários Idosos na Internet	18
Análise de Técnicas para Remoção de Fundo nas Imagens do Dataset BDICAFE	19
Análise do Cenário Competitivo e Educacional da Computação Quântica	20
Aprendizado baseado em Vision Transformer para Detecção dos Estádios de Maturação do Café	21
Aprimoramento da Técnica VOF-PLIC com Modelos de Aprendizado de Máquina	22
Aproximação Inteira da Distância Euclideana via Expansão Binária	23
Arquitetura FPGA baseada em Machine Learning para detecção de Ransomware	24
Avaliação da Swin-Unet para Segmentação de Lesões de Câncer de Endométrio em Imagens de Tomografia Computadorizada	25
Avaliação de Arquiteturas e Estratégias de Aumento de Dados para Classificação Multiclasse de Imagens Histológicas de Lesões no Pulmão	26
Avaliação de Efeito do Fine-Tuning em Modelos CNNs para a Classificação de Lesões Histológicas	27
Avaliação dos Modelos EfficientNetV2 e MobileNetV2 na Classificação de Imagens Histológicas de Displasia da Cavidade Oral	28
Avaliando Diferentes Metodologias para o Ensino de Computação Gráfica no Nível Superior ..	29
Calibração Evolutiva de um Modelo de Autômatos Celulares e Otimização no Posicionamento de Sinalização	30
Capacitando Profissionais de Saúde e Educadores no Diagnóstico de Doenças Raras e na Interpretação Genômica por meio da Bioinformática	31
Classificação Hierárquica Aplicada a Sistemas de Detecção de Intrusões em Redes IoT	32
Computação Inclusiva para Crianças e Jovens Neurodivergentes	33
DataRU: Ciência de dados aplicada aos Restaurantes Universitários - Resultados Finais	34
Deep Learning em Patologia Digital: Avaliação de ResNet e MobileNet na Análise Histológica de Displasia Oral e Pulmão	35

Dependabilidade do Sistema Operacional na Otimização de Modelos de Inteligência Artificial: Um Estudo de Caso Aplicado à Cibersegurança	36
Desafios no Pré-processamento de Dados Clínicos para Aplicação em Inteligência Artificial: Um Estudo de Caso em Reumatologia	37
Deteção de Diabetes Mellitus a partir da Espectroscopia de Infravermelho com Transformada de Fourier via Aprendizado de Máquina	38
EcoTech UFU: Tecnologia e Sustentabilidade na Destinação do resíduo eletrônico	39
Estudo da MobileNetV2 para Classificação de Câncer de Mama	40
Estudo de Desempenho de Modelos de Machine Learning na Identificação de Ataques do tipo Port Scanning	41
Estudo de técnicas para auxiliar no diagnóstico de Câncer de Mama	42
Evolução de um Sistema para Análise de Planos de Ensino baseado em Inteligência Artificial Generativa	43
GenPPI: Análise Comparativa de Interações Proteína-Proteína via Aprendizado de Máquina em Quatro Cepas de <i>Kosakonia cowanii</i> isoladas no Triângulo Mineiro	44
Geração Automática de Grades Horárias com Restrições de Disponibilidade para um Sistema de Distribuição de Disciplinas	45
GPT Teacher: Desenvolvimento de um Agente de LLM para Programação Assistida em Ambiente VSCode	46
Implementação e comparação de modelos ML na predição de cepas da Sars-Cov-2	47
Inspeção Heurística de uma Plataforma de Autoria de Jogos Educacionais	48
Integração de Memória Virtual Simulada à Plataforma SimulateOS	49
JuriZap: Um Assistente Jurídico Brasileiro Gratuito para Leigos baseado em Tecnologias de IA Generativa	50
Mapeamento Midiático do Viés Racial em Sistemas de Reconhecimento Facial: Resultados Preliminares	51
Método Automatizado para Deteção e Mitigação de Ataques Prompt Hacking em LLMs e SLMs	52
Modelo Neuro-Reconfigurável baseado em Modelagem Matemática para Deteção de Ataques DDoS com Baixa Latência	53
Modelos Residuais para Classificação de Displasia Oral e Lesões Pulmonares: Um Estudo Comparativo	54
O Impacto do Algoritmo de Grover na Segurança da Criptografia Simétrica	55
Otimização de Consultas em Banco de Dados Astronômico Massivo	56
Petúcia e a IA: Estratégia Prática para Desmistificar a Tecnologia na Educação Básica	57
Plataforma DebugandoED Repaginada: Integração de Gamificação e Dashboards no Aprendizado de Estruturas de Dados	58
Plataforma Integrada para Deteção e Análise de Ameaças Cibernéticas	59
Proposta de Implantação de um Chatbot para Avaliação no impacto das Operações do HC-UFU/EBSERH	60

Proposta de Intervenções Automatizadas via Chatbot na Plataforma DebugandoED	61
Recomendação de Formação de Equipes de Avaliação para o PNLD com Foco em Equidade e Representatividade	62
Reconhecimento de Emoções em Sinais Sonoros: Uma Revisão de Modelos e Metodologias	63
Redesign Centrado no Usuário: Uma Intervenção no Site das Bibliotecas da UFU sob a Perspectiva de Interação Humano Computador	64
Robustez ao Aumento de Dados: Impacto do Aumento de Dados nas Redes VGG e Squeeze-Net na Classificação de Lesões Histológicas	65
Seleção de Vértices Semiautomática e Reconstrução Automática de Malhas de Cabeças 3D para Fins de Animação	66
Sistema de Comparação de Preços para E-Commerce com Machine Learning	67
Sistema de Recomendação baseado em Conteúdo para Programação Competitiva: Uma Abordagem de Personalização de Treino	68
Tecnologia Acessível: Uso de VANTs de Baixo Custo no Monitoramento Ambiental e Ocorrências Críticas	69
Um Estudo Comparativo entre Sistemas de Gerenciamento de Bancos de Dados de Vetores com Foco em Dados Textuais	70
Um Sistema Tutor Inteligente para Apoio ao Ensino de Computação	71
Uma Abordagem Baseada em Transformer para Classificação de Espectros FTIR no Diagnóstico de Câncer Oral	72
Uma Ferramenta Computacional para Geração Automatizada de Perfis Geológicos a partir de Dados KML	73
Uma investigação do Algoritmo Genético aplicado ao Sudoku clássico 9×9	74
Uma solução baseada em Modelos de Linguagem e Geração Aumentada de Recuperação para auxílio ao discente da disciplina de Sistemas Operacionais	75
Uso de Heurísticas na Detecção de Mixer na Blockchain Bitcoin	76
Utilização de Redes Neurais com Otimização de Processamento de Imagens para detecção de Deficiências Nutricionais em Folhas de Café	77
2. Trilha: Trabalhos de Pós-graduação	78
Algoritmos Evolutivos para a Otimização Dinâmica de um Problema Discreto com Muitos Objetivos	79
Análise e Detecção de Anomalias de Cobertura em Redes Celulares: Uma Abordagem com Aumento de Dados e Modelagem Preditiva	80
Arquitetura Híbrida para Detecção fim a fim de Exfiltração DNS com Inteligência em Camada de Rede e Inspeção em Nuvem para Análise do Tráfego Criptografado	81
Avaliação de Fertilidade Seminal: Uma Abordagem Centrada na Análise dos Parâmetros de Morfologia, Vitalidade e Contagem Espermáticas com Base em Imagens Digitais e Deep Convolutional Neural Networks	82
Avaliação do Impacto de Metodologias Ativas e Pensamento Computacional no Desempenho em Programação Competitiva	83

Avaliações de Modelos de Deep Learning na Classificação Temporal de Deficiências Nutricio- nais em Folhas de Café	84
Cenários de Aplicação dos Knowledge Graphs em Sistemas de Recomendação: Uma Revisão Sistemática da Literatura	85
Classificação de Ameloblastoma e Carcinoma Ameloblástico em Patologia Digital: Uma Abor- dagem Eficiente com Modelos Híbridos Leves	86
Entre Rimas e Algoritmos: Uma Investigação sobre Tradução Automática Poética	87
Especialização Supervisionada de Modelos de Linguagem para Decisões Estratégicas no Pô- quer	88
Estudo de Destilação de Conhecimento para Classificação de Carcinoma Oral de Células Esca- mosas em Imagens Histológicas	89
IA na Educação: Caminhos para uma Transformação Pedagógica Centrada no Ser Humano ..	90
Impacto da Normalização e Redução de Dimensionalidade na Classificação de Imagens Histo- lógicas por Aprendizado de Máquina	91
Lógica Informal e Pensamento Crítico no Ensino Introdutório de Programação: Proposição e Validação de Modelo Pedagógico em contextos técnicos e de graduação	92
Modelagem Sequencial e Tabular do Crescimento e Produção em Plantios de Eucalipto	93
Modelo de Propagação da COVID-19 baseado em uma Rede de Autômatos Celulares	94
Novo Modelo de Simulação de Dinâmica de Pedestres para Ambientes Externos baseado em Autômatos Celulares	95
Perfil dos Participantes em Treinamentos de Programação Competitiva em uma Instituição de Ensino Superior Moçambicano (2024–2025)	96
Predição de Desempenho Acadêmico Discente Utilizando Dados de Frequência em Aulas: Uma Abordagem com Múltiplos Modelos de Aprendizado de Máquina	97
Projeto e Verificação da Olimpíada Brasileira de Programação Quântica	98
Técnicas de IA Explicável aplicadas ao Problema de Manutenção Preditiva na Indústria Ali- mentícia	99
Um panorama da Infraestrutura Tecnológica para Implementação da Programação Competiti- va em Moçambique	100
Uma Nova Arquitetura Híbrida para Otimização Multiobjetivo de Hiperparâmetros em IDS ...	101
Visualizações em Vídeos de Vigilância Orientadas por Ontologia	102

Trilha: Trabalhos de Graduação

Aleatoriedade Quântica em Sistemas Fotônicos: Da Simulação Computacional à Criptografia Moderna

Evellyn Fernanda Machado, Giulia Rodrigues

Universidade Federal de Uberlândia (UFU)

{evellyn.machado, giullia.rodrigues}@ufu.br

Introdução: A geração de números verdadeiramente aleatórios é essencial para sistemas modernos de segurança, sobretudo em aplicações de criptografia avançada e cenários pós-quânticos. Geradores pseudoaleatórios não garantem imprevisibilidade absoluta, e métodos quânticos baseados em qubits ainda enfrentam limitações de velocidade, eficiência e custo. Nesse contexto, sistemas fotônicos de variáveis contínuas (CV) surgem como alternativa promissora, explorando as flutuações quânticas do vácuo como fonte de aleatoriedade física. Essa abordagem utiliza tecnologias ópticas maduras, alta taxa de detecção e compatibilidade com plataformas de telecomunicações, permitindo taxas superiores de geração e maior robustez operacional. Este trabalho apresenta uma simulação prática de um gerador quântico de números aleatórios baseado em sistemas fotônicos CV, demonstrando o processo de geração e avaliando seu potencial para uso em criptografia. **Objetivos:** Demonstrar, por meio de simulação, que a abordagem fotônica de Variáveis Contínuas (CV) é mais eficiente e prática para geração de aleatoriedade em alta velocidade. A análise justifica-se pela necessidade de alinhar o desenvolvimento de QRNGs às demandas reais de hardware (Gbps, robustez e custo), aspecto ainda pouco abordado em métodos DV. Secundariamente, busca-se validar a aplicabilidade dos bits aleatórios gerados em um protocolo criptográfico padrão. **Método:** A simulação é realizada em Python, utilizando a biblioteca Strawberry Fields, modelando um circuito fotônico CV que mede quadraturas do campo eletromagnético. O processo simula a detecção das flutuações quânticas do vácuo por medições de quadratura (X ou P), resultando em valores contínuos associados à imprevisibilidade quântica. Esses valores são então processados para extração de bits aleatórios, posteriormente usados como chave em uma cifra XOR, demonstrando sua aplicação criptográfica. **Resultados:** Os resultados esperados incluem a geração de dados com distribuição estatística compatível com ruído quântico, evidenciando a capacidade do modelo fotônico CV de produzir valores aleatórios imprevisíveis. A aplicação dos dados em uma cifra XOR valida sua funcionalidade criptográfica, mostrando que os bits gerados cumprem o papel de chave segura. Além disso, a simulação reforça as vantagens dos sistemas fotônicos CV, como escalabilidade, operação com componentes ópticos acessíveis e taxas de geração superiores às abordagens baseadas em qubits. **Conclusão:** A simulação demonstra que sistemas fotônicos de variáveis contínuas são uma alternativa viável e eficiente para geração quântica de aleatoriedade, oferecendo vantagens em velocidade, robustez e implementação prática. Embora executada em ambiente clássico, a abordagem ilustra o funcionamento físico do método e destaca seu potencial para aplicações criptográficas modernas, contribuindo para o avanço de soluções quânticas acessíveis e escaláveis em segurança da informação.

Trilha: Trabalho de Graduação.

Algoritmo Genético Orientado à Desordem para a Otimização da Exploração de Ambientes por Enxames de Robôs

Gabriel A. C. Batista, Claudiney R. Tinoco

Universidade Federal de Uberlândia (UFU)
{gabriel.alcantara, claudiney.tinoco}@ufu.br

Introdução: A coordenação de enxames de robôs representa um desafio significativo na robótica contemporânea, particularmente em tarefas de cobertura ambiental onde a distribuição espacial equilibrada é crucial. Este trabalho investiga a aplicação de Algoritmos Genéticos (AG) para otimização automática dos parâmetros do modelo PheroCom de coordenação de enxames, com foco específico no aprimoramento da homogeneidade da exploração do ambiente. Propõe-se a integração da entropia como métrica fundamental na função de aptidão do AG, visando superar as limitações de configurações manuais e abordar o problema de distribuição espacial desbalanceada. **Objetivos:** Desenvolver um método de otimização automática dos parâmetros do modelo PheroCom de coordenação de enxames, utilizando Algoritmos Genéticos com função de aptidão baseada em entropia para melhorar a homogeneidade da cobertura ambiental. **Método:** Implementou-se um Algoritmo Genético com função de fitness personalizada, onde a entropia atua como métrica de avaliação da distribuição espacial dos robôs. A abordagem utiliza a entropia como índice multiplicativo em um sistema de penalização, permitindo identificar configurações que favorecem distribuições mais equilibradas no ambiente de simulação. **Resultados:** Os experimentos demonstraram que a integração da entropia permitiu uma análise eficaz da homogeneidade da exploração, resultando em uma ocupação 24,75% mais eficiente em comparação com configurações manuais. A métrica mostrou-se adequada para promover distribuições espaciais mais equilibradas, evitando tanto concentrações excessivas quanto regiões negligenciadas. **Conclusão:** Os resultados obtidos validam a utilidade da entropia como métrica para otimização da coordenação em enxames de robôs e demonstram a eficácia dos AGs no ajuste automático de parâmetros de modelos de coordenação. A abordagem proposta constitui uma contribuição significativa para o campo, oferecendo um método sistemático para melhorar a eficiência na cobertura ambiental. Trabalhos futuros incluem a incorporação de métricas complementares, exploração de técnicas de paralelização para redução do tempo computacional, e extensão do método para outros contextos aplicados como coordenação de UAVs em agricultura de precisão.

Trilha: Trabalho de Graduação.

Análise Bibliométrica da Inteligência Artificial na Agricultura: Mapeamento de Tendências e Frentes de Pesquisa

Mariana S. F. Souza, Claudiney R. Tinoco

Universidade Federal de Uberlândia (UFU)
{mariana.souza2, claudiney.tinoco}@ufu.br

Introdução: A aplicação de técnicas de Inteligência Artificial (IA) na agricultura tem ganhado crescente relevância acadêmica e prática, impulsionando transformações significativas no setor agrícola global. Diante do acelerado desenvolvimento tecnológico, torna-se fundamental compreender o panorama científico e identificar direções emergentes. Este estudo apresenta um mapeamento bibliométrico preliminar das publicações sobre IA aplicada à agricultura, fornecendo uma visão abrangente e sistematizada do estado da arte. A pesquisa busca não apenas catalogar as tendências existentes, mas também analisar criticamente a distribuição temática e geográfica do conhecimento, oferecendo insights valiosos para a comunidade. **Objetivos:** O trabalho tem como objetivo principal analisar e visualizar dados bibliométricos de dez temas representativos da IA aplicada à agricultura. Especificamente, busca-se: (i) identificar as áreas de concentração temática com maior volume de publicações; (ii) mapear a distribuição geográfica das pesquisas; e, (iii) comparar a maturidade tecnológica entre diferentes subáreas. **Método:** O procedimento metodológico seguiu três etapas sequenciais: (i) coleta sistemática na base IEEE Xplore utilizando queries específicas para cada tema; (ii) exportação e organização dos metadados em formato CSV; e (iii) processamento e análise dos dados mediante a ferramenta Power BI, permitindo a criação de painéis interativos e visualizações analíticas. **Resultados:** As análises realizadas permitiram a construção de um panorama quantitativo das publicações, revelando tendências temporais significativas e concentrações geográficas específicas. A comparação entre os temas evidenciou a predominância de determinadas subáreas de IA, enquanto outras apresentaram volume reduzido de produção científica, sinalizando temas emergentes e áreas que demandam maior investigação. **Conclusão:** Os resultados validam a relevância crescente da IA no domínio agrícola e demonstram a utilidade da análise bibliométrica como ferramenta de diagnóstico do estado da arte. O mapeamento realizado constitui uma base sólida para orientar futuras investigações, destacando fronteiras de inovação promissoras e direcionando esforços de pesquisa para tópicos com potencial de desenvolvimento. Estudos futuros poderão expandir a análise para outras bases de dados e incorporar métricas de impacto e colaboração científica.

Trilha: Trabalho de Graduação.

Análise de Comentários de Vídeos do YouTube utilizando Redes Textuais

Eduarda Moura, Fabíola Pereira

Universidade Federal de Uberlândia (UFU)
{eduarda.lobes, fabiola.pereira}@ufu.br

Introdução: Para além de nuvens de palavras, dados textuais podem ser analisados e explorados quando modelados em rede. A análise de redes textuais, com a representação de textos como grafos formados por nós e arestas, permite extrair informações valiosas sobre a estrutura e dinâmica textual. Embora existam diversas abordagens para construir redes a partir de dados textuais, a literatura carece de estudos comparativos acerca do impacto das diferentes estratégias de modelagem destas. Em especial, dados de redes sociais são utilizados, pois apresentam grande diversidade em suas características linguísticas e estruturais. Torna-se necessário compreender como diferentes estratégias de construção de redes textuais afetam a qualidade e utilidade das análises, especialmente em tarefas como mineração de opinião e análise exploratória. Os comentários do YouTube exemplificam esse contexto, combinando volume de dados e diversidade linguística, que são desafios contemporâneos para esse tipo de análise. **Objetivos:** O estudo visa explorar técnicas de representação de redes textuais de maneira comparativa, avaliando diferentes estratégias de modelagem e suas respectivas métricas de análise. A partir desse objetivo, a avaliação de desempenho das redes modeladas em tarefas de mineração de opinião precede o cerne da pesquisa, que consiste nas análises exploratórias baseadas em métricas específicas, comparando-se os resultados obtidos. Espera-se elaborar diretrizes metodológicas para orientar a escolha de técnicas de modelagem de redes textuais em contextos similares à análise de dados sociais online. **Método:** Inicialmente, realiza-se a coleta e pré-processamento dos dados, preparando-os para serem utilizados na construção de redes textuais. Para a fase inicial, utilizou-se um corpus de aproximadamente 53.000 comentários. Em seguida, serão realizadas análises estruturais das redes geradas, utilizando métricas de grafos e técnicas de detecção de comunidades e identificação de termos centrais, dividindo-se em análise exploratória e mineração de opinião. Por fim, os resultados serão sistematizados para propor diretrizes comparativas que orientem futuras aplicações de redes textuais em dados de redes sociais. **Resultados:** Até o momento, a revisão bibliográfica tem sido o foco do projeto, proporcionando uma compreensão aprofundada sobre redes textuais temporais, ciência de rede e mineração de opinião. Esse estudo possibilitou um esboço inicial de uma rede textual a partir da base de dados coletada, com destaque nos nós com maiores valores de centralidade. **Conclusão:** A expectativa é fornecer contribuições metodológicas significativas para análise de redes textuais, demonstrando que diferentes modelagens possuem vantagens e limitações específicas para tarefas distintas. Espera-se que as diretrizes propostas auxiliem pesquisadores na seleção de técnicas de modelagem, contribuindo para maior rigor metodológico em pesquisas de mineração de texto e análise de redes sociais.

Trilha: Trabalho de Graduação.

Análise de Diferentes Operações de Aumento de Dados em Modelos de Classificação de Imagens Histológicas

João Antônio M. Jordão, Marcelo Nascimento

Universidade Federal de Uberlândia (UFU)
{joao.jordao, marcelo.nascimento}@ufu.br

Introdução: Nos últimos anos, a utilização de visão computacional no campo da saúde vem crescendo com o desenvolvimento de sistemas cada vez mais precisos e confiáveis, como exemplo as técnicas aplicadas na análise de tecidos histológicos. Em paralelo, o número limitado de dados e a morosidade do processo de montagem de um dataset com imagens histológicas demonstram um gargalo no desenvolvimento dessas tecnologias. Considerando esse quadro, a técnica de aumento de dados (utilização de dados existentes para criação de novas amostras) se mostra uma estratégia muito útil, auxiliando na construção desse tipo de dataset. Nesse contexto, a análise dos impactos das diferentes operações de aumento de dados em modelos de classificação de imagens histológicas confere um maior entendimento dessa técnica e permite o refinamento dos resultados desses modelos. **Objetivos:** Estudar as relações entre a aplicação de diferentes operações de aumento de dados em um dataset de imagens histológicas e os resultados de um modelo de aprendizagem profunda para análise e classificação de displasias na cavidade oral, buscando otimizá-lo. **Método:** Foi utilizada uma rede neural residual convolucional (modelo ResNet-50) para produzir um classificador capaz de reconhecer e classificar níveis de displasia juntamente com um algoritmo de aumento de dados configurável, permitindo a seleção de operações (desde simples manipulações de imagens até aplicação de transformações e distorções mais complexas) e tamanhos específicos de aumento. O dataset aumentado apresenta imagens histológicas da cavidade oral com diferentes níveis de displasias, com inicialmente 114 imagens para cada uma das seguintes classes: saudável (não apresenta displasia), leve, moderada e severa. **Resultados:** Mesmo com o trabalho em etapa de finalização, a abordagem proposta se mostrou capaz de demonstrar as relações entre as diferentes operações de aumento de dados e os resultados do modelo de análise e classificação construído, permitindo a análise do impacto dessas diferentes operações em imagens histológicas e a otimização do modelo de classificação de displasias, que atingiu uma acurácia de 99,66%. **Conclusão:** O estudo da combinação de tecnologias como redes neurais convolucionais e estratégias de aumento de dados para confecção de um modelo de visão computacional permite um maior entendimento dessas interações e a construção de modelos mais robustos, contribuindo para o desenvolvimento de ferramentas de visão computacional contextualizadas com imagens histológicas e tecnologias orientadas à aplicação médica.

Trilha: Trabalho de Graduação.

Análise de Resolvedores de DNS para Navegação Segura de Usuários Idosos na Internet

Jaqueline G. Souza, Thiago Ribeiro

Universidade Federal de Uberlândia (UFU)
{jaqueline.goncalves, tpribeiro}@ufu.br

Introdução: Este trabalho investigará o uso do DNS seguro como estratégia de proteção digital para usuários leigos e pessoas idosas. O crescente uso da Internet e a vulnerabilidade desse grupo diante de ameaças virtuais tornam essencial a análise de soluções que ofereçam segurança e acessibilidade. Nesse contexto, serão examinados quatro serviços de resolução segura de nomes de domínio disponíveis no mercado: Cloudflare DNS, NextDNS, AdGuard DNS e Quad9, considerando suas funcionalidades e formas de filtragem de conteúdo para proteção digital desse público. **Objetivos:** O objetivo deste estudo é avaliar a eficácia dos resolvedores de DNS, mencionados anteriormente, na detecção e no bloqueio de sites potencialmente perigosos, considerando sua aplicabilidade como ferramenta de proteção digital voltada a pessoas idosas. Busca-se comparar os mecanismos de filtragem empregados por cada serviço, analisar sua eficiência na prevenção de acessos a domínios maliciosos e examinar a facilidade de configuração oferecidas, levando em conta as limitações tecnológicas frequentemente observadas nesse público. **Método:** A pesquisa adota abordagem exploratória, com realização de testes práticos dos resolvedores mencionados. Cada serviço será avaliado em cenários simulados que reproduzam situações típicas de navegação usuários leigos e pessoas idosas, incluindo tentativas de acesso a sites potencialmente nocivos. **Resultados:** Espera-se que a análise comprove a efetividade do uso de resolvedores de DNS como uma medida prática de segurança para usuários leigos e pessoas idosas, evidenciando que há serviços de DNS que se destacam pela capacidade de bloqueio e pela facilidade de uso. A comparação prática entre os resolvedores permitirá identificar quais apresentam melhor equilíbrio entre proteção, desempenho e acessibilidade, demonstrando que essa tecnologia pode oferecer uma camada adicional de segurança mesmo para usuários com pouca familiaridade ou malícia digital. **Conclusão:** O estudo demonstrará que o uso de DNS seguro pode representar uma alternativa viável e acessível de proteção digital para usuários leigos e pessoas idosas, desde que acompanhado de mecanismos de segurança confiáveis. De forma complementar, para idosos e usuários com maior familiaridade tecnológica, é possível combinar esses resolvedores com ferramentas adicionais, como o pi-hole, ampliando as camadas de filtragem e controle de conteúdo e garantindo uma proteção ainda mais completa. A análise comparativa dos serviços contribuirá para compreender a relevância dessa tecnologia como instrumento de segurança online e para orientar futuras ações voltadas a proteção desse público.

Trilha: Trabalho de Graduação.

Análise de Técnicas para Remoção de Fundo nas Imagens do Dataset BDICAFE

Gabriel Tassara, Thiago Ribeiro

Universidade Federal de Uberlândia (UFU)

{tassara, tpribeiro}@ufu.br

Introdução: Esta pesquisa surgiu da necessidade de otimizar os resultados do projeto “Identificação precoce de deficiência nutricional em cafeeiros utilizando descritores digitais e técnicas de aprendizado profundo”. O projeto original utiliza descritores para processar imagens e gerar arrays por meio da biblioteca Numpy, treinando os modelos de IA de forma manual. A partir disso, surgiu a pergunta: “A remoção do fundo das imagens do dataset BDICAFE resultaria em ganhos na validação do modelo?”. **Objetivos:** Um dos objetivos deste trabalho é a remoção do fundo das imagens do dataset BDICAFE de forma automatizada, eficiente e sem perdas significativas no objeto alvo da imagem (a folha de café). **Método:** Foram realizadas comparações preliminares com 6 modelos, sendo que dois deles foram imediatamente descartados por apresentarem resultados insatisfatórios, não conseguindo manter a estrutura completa da folha após a remoção do fundo. Foram selecionados 4 modelos para a próxima fase de análise mais aprofundadas: U2NET, SAM, ISNET-GENERAL-USE (todos da biblioteca REMBG) e o InsSPyReNet. Para estabelecer o padrão-ouro (Ground Truth), foram realizados recortes manuais de folhas de referência de “forma cirúrgica”. As máscaras geradas pelos modelos foram então, comparadas com este padrão-ouro utilizando a função `jaccard_score` (da biblioteca `sklearn.metrics`). Essa métrica calcula a “Interseção sobre a União” (IoU), que mede o grau de similaridade entre as máscaras geradas pelos modelos e a máscara ideal. Além da acurácia, um fator importante a se considerar é a eficiência. Nesse quesito, o U2NET demonstrou maior velocidade com menor uso de recursos, em contraste com o SAM (desenvolvido pela empresa META), que se mostrou bastante exigente, apesar de seus resultados satisfatórios. **Resultados:** Os resultados parciais indicam que o modelo ISNET-GENERAL-USE apresentou maior semelhança com o padrão-ouro, porém devido ao número limitado de testes, este resultado não pode ser conclusivo e, até o momento não se consegue afirmar a eficiência de todos os métodos. **Conclusão:** A pesquisa continuará a buscar o equilíbrio entre acurácia e eficiência (qualidade/tempo), visto que o projeto principal utiliza o dataset BDICAFE com mais de 16.000 imagens em resolução 4K e os recursos computacionais dos pesquisadores são limitados. Espera-se que este projeto consiga obter a melhor solução em acurácia e eficiência para auxiliar no ganho de performance do projeto original.

Trilha: Trabalho de Graduação.

Análise do Cenário Competitivo e Educacional da Computação Quântica

Ellen Amaral Santana, Giullia Rodrigues, João Henrique Souza

Universidade Federal de Uberlândia (UFU)

{ellen.christina, giullia.rodrigues, joaohs}@ufu.br

Introdução: O desenvolvimento exponencial da Computação Quântica (CQ) gera alta demanda por profissionais qualificados. O trabalho analisa a formação do ecossistema educacional e competitivo de CQ no Brasil, comparando sua maturidade com a programação clássica. **Objetivos:** O foco é mapear cursos e competições de CQ para identificar necessidades e o estágio de desenvolvimento nacional na área. **Método:** A pesquisa combinou pesquisa bibliográfica e documental. As unidades de análise foram categorizadas em três grupos: competições de programação quântica (com simuladores quânticos), competições de programação clássica (para comparação) e a oferta de cursos de CQ no Brasil. Os dados coletados foram submetidos a uma análise comparativa detalhada. **Resultados:** Constatou-se uma acentuada disparidade de maturidade entre os formatos de competição (clássicas desde 1989, quânticas a partir de 2020). As competições clássicas contam com um modelo maduro, predominantemente presencial (85,7%), individual (71,4%), de curta duração (1-2 dias) com forte presença na educação básica (53,3%). As competições quânticas possuem um modelo emergente, majoritariamente remoto (62,5%), em equipe (55,6%), de longa duração (meses/anos) com público-alvo mais avançado (40% superior/50% geral). Em relação aos cursos de CQ no Brasil, a oferta é limitada (apenas 12 opções), majoritariamente de nível profissionalizante (50%) ou especialização (33%), no formato remoto (58,3%) e de curta duração (até 60h em 46,2% dos casos). O público-alvo principal é formado por estudantes de graduação (44,4%) e profissionais (27,8%). Os custos estão divididos: 37,5% gratuitos e 37,5% de alto custo (>R\$1.000). **Conclusão:** O formato flexível das competições de CQ (remoto, em equipe, longa duração) é uma vantagem para impulsionar a inovação e atrair talentos avançados, ao contrário das competições de programação clássicas que focam na avaliação pontual de habilidades. Contudo, percebemos que existe limitação na oferta educacional brasileira, que é restrita e concentrada em capacitação rápida e direcionada ao mercado. Para consolidar a área no país, a infraestrutura educacional precisa expandir-se e diversificar-se. O trabalho fornece um mapeamento detalhado, sendo base para futuras políticas de ensino e investimento em tecnologia quântica.

Trilha: Trabalho de Graduação.

Aprendizado baseado em Vision Transformer para Detecção dos Estádios de Maturação do Café

Wallace G. M. Pinheiro, Murillo G. Carneiro, Cleyton B. Alvarenga

Universidade Federal de Uberlândia (UFU)

{wallacegmp, mgcarneiro, cleytonalvarenga}@ufu.br

Introdução: Este trabalho está inserido no contexto do Aprendizado Profundo aplicado à agricultura, especificamente à cultura do café. Com o avanço da Agricultura 4.0 e o fortalecimento da agricultura de precisão, técnicas de Visão Computacional têm sido utilizadas para auxiliar produtores em etapas críticas, como determinar o momento ideal de colheita do fruto do café. Nesse cenário, abordagens baseadas em Vision Transformers (ViTs) têm se destacado por capturar relações globais nas imagens, oferecendo análises mais precisas e robustas. Assim, este estudo investiga o uso de ViTs na caracterização dos estádios de maturação dos grãos de café e na decisão de colher ou não o fruto, visando maior eficiência e qualidade na colheita. **Objetivos:** Avaliar o desempenho de um modelo baseado em Vision Transformers na classificação de imagens de frutos de café, analisando sua capacidade de identificar os diferentes estádios de maturação e apoiar a decisão de colheita. A escolha do ViT se justifica por seu desempenho superior em tarefas de classificação de imagens complexas, superando limitações de arquiteturas baseadas em convoluções. **Método:** O estudo utilizou uma base de dados com imagens de frutos de café em cinco estádios distintos de maturação. Essas classes foram usadas na classificação, enquanto, para a decisão de colheita, consideraram-se duas classes: “colher” e “não colher”. As imagens passaram por pré-processamento e data augmentation com rotação, zoom e espelhamento, para aumentar a variabilidade e melhorar o desempenho do modelo. O Vision Transformer foi ajustado por fine-tuning, com pesos pré-treinados, aprimorando a generalização e a precisão. Embora modelos como a MobileNet apresentem bons resultados, o ViT se destaca por capturar relações globais na imagem, melhorando a identificação dos estádios de maturação. **Resultados:** Foram utilizadas cerca de 5.750 imagens para treino, 1.700 para validação e 500 para teste. Na identificação dos estádios de maturação, com cinco classes, o modelo alcançou acurácia de 86,40%, com precisão entre 0,73 e 1,00, recall entre 0,66 e 0,98 e F1-score de 0,76 a 0,94. Na decisão de colheita, com duas classes (colher e não colher), obteve acurácia de 97,60% e F1-score de até 0,98. Em comparação com abordagens anteriores baseadas em redes convolucionais, o Vision Transformer apresentou desempenho superior, evidenciando sua capacidade de capturar relações globais e aprimorar a precisão em cenários agrícolas. **Conclusão:** Os resultados demonstram a eficiência do modelo ViT na classificação dos estádios de maturação e na decisão de colheita. As acurácias de 86,40% e 97,60% confirmam o bom desempenho e a consistência do modelo. Conclui-se que o Vision Transformer é uma alternativa promissora para aplicações agrícolas baseadas em visão computacional, contribuindo para a automação e precisão no processo de colheita do fruto do café.

Trilha: Trabalho de Graduação.

Aprimoramento da Técnica VOF-PLIC com Modelos de Aprendizado de Máquina

Vinicius Martins, Rafael Figueiredo

Universidade Federal de Uberlândia (UFU)
{vinicius.tavares, rafigueiredo}@ufu.br

Introdução: A simulação numérica de escoamentos multifásicos é essencial em diversas aplicações da engenharia. Um dos principais desafios consiste em representar e atualizar a interface entre as fases com precisão. Nesse contexto, a técnica Volume-of-Fluid (VOF), combinada com a Piecewise Linear Interface Calculation (PLIC), destaca-se como uma abordagem computacional eficaz para a reconstrução de interfaces. O sucesso do método PLIC depende da estimativa precisa das normais da interface. Diante das limitações dos métodos tradicionais, a aplicação de técnicas de inteligência artificial (IA) surge como uma alternativa promissora para aprimorar a precisão da reconstrução. **Objetivos:** A reconstrução de interfaces de forma precisa e eficiente é um dos grandes desafios na simulação de escoamentos multifásicos. O avanço recente das técnicas de IA tem mostrado resultados expressivos em problemas complexos de regressão e previsão de campos físicos. Assim, o objetivo principal desta pesquisa é desenvolver métodos alternativos, baseados em aprendizado de máquina, para estimar de maneira eficiente os vetores normais em escoamentos bifásicos simulados pela técnica VOF-PLIC. **Método:** Para atingir o objetivo proposto, o primeiro passo consiste na construção de uma base de dados sólida e precisa. Para isso, realiza-se a reconstrução de interfaces simples utilizando a abordagem VOF. Com o campo VOF definido, calculam-se analiticamente os vetores normais, gerando dados altamente precisos. Essa base será então utilizada para o treinamento do modelo de IA, que busca aprender a relação entre o campo VOF e as normais correspondentes. Após o treinamento, os resultados obtidos pelo modelo serão avaliados quantitativamente e comparados com os produzidos por alguma técnica da literatura, em especial a técnica ELVIRA (Efficient Least Squares Volume-of-Fluid Interface Reconstruction Algorithm), de modo a verificar ganhos em precisão e custo computacional. **Resultados:** Os resultados obtidos até o momento incluem a implementação inicial do campo do VOF, bem como o cálculo analítico das normais para formas simples, com o objetivo de compor uma base de dados sólida e confiável para o treinamento do modelo de machine learning. Com esses dados, espera-se que o modelo seja capaz de estimar vetores normais com rapidez e precisão, inclusive para interfaces complexas, apresentando desempenho comparável ou superior à técnica ELVIRA. **Conclusão:** A pesquisa apresenta caráter interdisciplinar, envolvendo matemática aplicada, engenharia e ciência da computação. Trata-se de um projeto promissor que, uma vez concluído com êxito, poderá contribuir significativamente para a simulação numérica de escoamentos multifásicos e, com o uso de ferramentas modernas, aprimorar de forma substancial a abordagem computacional VOF-PLIC, oferecendo uma alternativa eficiente às técnicas tradicionais de reconstrução de interface.

Trilha: Trabalho de Graduação.

Aproximação Inteira da Distância Euclidiana via Expansão Binária

Gustavo Luis Nascimento

Universidade Federal de Uberlândia (UFU)

gustavo.luis@ufu.br

Introdução: A pesquisa objetiva examinar e criar versões aprimoradas de algoritmos que usam operações com números inteiros. A computação de distâncias N-dimensionais é relevante em mineração de dados, mas a distância euclidiana tradicional usa ponto flutuante, que é computacionalmente custoso. O projeto foca em adaptar esse cálculo para operações inteiras, usando a expansão binária como aproximação de raiz quadrada. Busca-se simplificar a implementação em hardware, como FPGAs, para obter paralelização e melhoria no tempo de execução. **Objetivos:** O objetivo geral é elaborar uma versão do cálculo de distância N-D com operações inteiras e sua implementação em hardware. Os objetivos específicos são: 1) Realizar um levantamento do estado da arte sobre otimização de algoritmos para hardware; 2) Reestruturar algebricamente a equação da distância para eliminar operações com números reais; 3) Desenvolver um protótipo em software; 4) Executar estudos comparativos de precisão e desempenho; 5) Desenvolver uma versão em VHDL, sintetizá-la em FPGA e avaliar o desempenho em hardware. **Método:** A metodologia é sequencial. Após a Revisão Bibliográfica, será feita a Reestruturação da Equação da Distância, manipulando a fórmula para eliminar a raiz quadrada e reais usando a expansão binária (ex: $A + B/2$). Um Protótipo em Software será criado para validar a lógica e precisão. Segue-se um Estudo Comparativo de precisão e desempenho com a abordagem euclidiana. O algoritmo validado será descrito em VHDL, sintetizado em FPGA e submetido a testes práticos na Avaliação de Desempenho em Hardware. **Resultados:** Os resultados esperados incluem uma versão integralizada e funcional do algoritmo de distância, materializada em um protótipo de software validado e uma implementação VHDL. O principal resultado será o estudo comparativo, quantificando os benefícios de desempenho (tempo de execução) e o custo de precisão (erro médio absoluto) da abordagem inteira, discutindo-os com a literatura. Espera-se que a versão em hardware demonstre melhorias significativas no tempo de execução. **Conclusão:** Espera-se concluir que a adaptação do algoritmo para operações inteiras, via expansão binária, é uma abordagem eficiente e viável. A pesquisa deverá validar que a nova formulação, embora seja uma aproximação que pode ter erros de precisão, mantém boa precisão e melhora o desempenho computacional, viabilizando a implementação em hardware digital. A discussão final indicará a relevância da solução, suas vantagens em FPGAs e suas limitações.

Trilha: Trabalho de Graduação.

Arquitetura FPGA baseada em Machine Learning para detecção de Ransomware

Camilly Cruvinel, Diego Molinos

Universidade Federal de Uberlândia (UFU)
{camilly.cruvinel, diego.molinos}@ufu.br

Introdução: O crescimento de ataques de ransomware tem representado uma das maiores ameaças à segurança da informação, afetando não apenas empresas, mas também infraestruturas críticas. O principal desafio associado a esse tipo de malware reside no fato de que sua detecção ocorre antes do início do processo de cifragem dos dados, etapa que, uma vez concluída, torna a recuperação dos dados significativamente complexa. Embora técnicas de Machine Learning (ML) tenham sido aplicadas com sucesso na detecção deste tipo de ameaça, os modelos produzidos ainda enfrentam limitações quando projetados em ambiente real, sendo uma das principais delas a dependência de frameworks e do sistema operacional, que introduzem latência e aumentam o tempo de inferência. **Objetivos:** Diante desse cenário, o objetivo deste trabalho é desenvolver uma arquitetura reconfigurável de detecção acelerada baseada em FPGA (Field-Programmable Gate Array), reduzindo a latência, aumentando o desempenho e viabilizando a detecção de ransomware ainda nas fases iniciais de execução. **Método:** Metodologicamente, esta pesquisa é de caráter experimental e exploratório e, encontra-se estruturada em quatro etapas principais, são elas: (a) estudo do estado da arte sobre detecção de ransomware e aceleração de inferência em hardware; (b) treinamento de um classificador do tipo Árvore de Decisão utilizando o framework Scikit-learn/Python e o dataset (CIC-MalMem-2022); (c) implementação de uma aplicação em linguagem C que materializa a árvore de decisão exportada do classificador, etapa importante para garantir a interpretabilidade do modelo de árvore gerado; (d) implementação de uma arquitetura reconfigurável em FPGA a partir da estrutura da árvore, eliminando a dependência de frameworks e do sistema operacional, reduzindo a latência de inferência. A arquitetura em FPGA desenvolvida foi avaliada seguindo os mesmos parâmetros dos testes das implementações em Python e em linguagem C. **Resultados:** Após os testes, a arquitetura proposta apresentou redução significativa no tempo de inferência, alcançando ganho de desempenho superior a 150% em relação às implementações equivalentes em linguagem C e Python, mantendo os mesmos critérios de classificação e a acurácia do modelo original. **Conclusão:** Esses resultados evidenciaram o potencial do uso de FPGAs para detecção comportamental de ransomware em tempo quase real, reforçando sua viabilidade em sistemas que demandam alta responsividade e elevada confiabilidade.

Trilha: Trabalho de Graduação.

Avaliação da Swin-Unet para Segmentação de Lesões de Câncer de Endométrio em Imagens de Tomografia Computadorizada

Bruno Tomé, Fernanda Santos

Universidade Federal de Uberlândia (UFU)
{brunoferreira.tome, fmcsantos}@ufu.br

Introdução: O Câncer de Endométrio (CE) é um dos tumores ginecológicos mais frequentes, e seu diagnóstico precoce continua sendo um desafio, principalmente devido à subjetividade envolvida na interpretação de imagens médicas. O uso de Diagnóstico Auxiliado por Computador surge como uma alternativa promissora, mas limitado pela falta de conjunto de dados específicos para a doença. A publicação do dataset ECPC-IDS trouxe o primeiro benchmark em grande escala voltado para essa tarefa. Entretanto, as análises nesse conjunto de dados apontaram um novo desafio, a segmentação das lesões ainda apresenta baixo desempenho. Ou seja, as arquiteturas tradicionais de Redes Neurais Convolucionais (CNNs), tem dificuldade em identificar os limites dos tumores nesse tipo de imagem, evidenciando a necessidade de novas abordagens mais sensíveis a variações de contraste e textura.

Objetivos: O principal objetivo é implementar e avaliar a Swin-Unet para enfrentar o desafio da segmentação de imagens de Tomografia Computadorizada (CT) presentes no dataset ECPC-IDS. Além disso, busca-se comparar seu desempenho com o de arquiteturas baseadas em CNNs, como a U-Net com backbones VGG-16 e ResNet50, para verificar se os mecanismos de atenção da Swin-Unet conseguem melhorar a segmentação das lesões de CE, neste tipo de exame.

Método: O conjunto de dados utilizado foi definido por um subconjunto criado a partir de 1516 imagens CT e suas respectivas máscaras de ground-truth do dataset ECPC-IDS. Inicialmente, aplicou-se um pipeline de pré-processamento padronizado, que incluiu o redimensionamento das imagens para 224×224 pixels, conversão para escala de cinza e a normalização dos valores de pixel para o intervalo $[0, 1]$. O conjunto de dados foi dividido em três partes: 70% para treino, 15% para validação e 15% para teste. No conjunto de treino, foram aplicadas técnicas de aumento de dados, como transformações elásticas, rotações e inversões horizontais, para ampliar a variabilidade das amostras e reduzir o risco de overfitting. Foram comparadas duas abordagens, o modelo Swin-Unet e os modelos U-Net com backbones VGG16 e ResNet50, as quais utilizaram aprendizado por transferência com pesos pré-treinados na ImageNet.

Resultados: A avaliação no conjunto de teste demonstrou que a otimização do batch size foi crucial. A Swin-Unet com batch size 16 alcançou o melhor desempenho, obtendo a menor perda de teste 0.0037 e a maior Sensibilidade 0.8467, superando os modelos U-Net. O modelo ResNet50 obteve a maior precisão 0.8779, mas uma Sensibilidade inferior 0.7942. O modelo VGG16 apresentou o desempenho mais baixo, com sensibilidade de 0.7284.

Conclusão: Conclui-se que os resultados validaram a hipótese de que a arquitetura Transformer é superior para esta tarefa de segmentação. Após a otimização de hiperparâmetros, a Swin-Unet foi o modelo mais robusto, superando as CNNs pré-treinadas para diagnóstico médico.

Trilha: Trabalho de Graduação.

Avaliação de Arquiteturas e Estratégias de Aumento de Dados para Classificação Multiclasse de Imagens Histológicas de Lesões no Pulmão

Gustavo Miranda¹, Daniel Gonçalves¹,
Domingos Oliveira², Marcelo Nascimento¹

¹Universidade Federal de Uberlândia (UFU)

²Instituto Federal de São Paulo (IFSP)

{gustavo.miranda, danielbgoncalves, marcelo.nascimento}@ufu.br
domingos.latorre@gmail.com

Introdução: O câncer de pulmão é uma das principais causas de morte por neoplasias, com cerca de 2,2 milhões de novos casos anuais. Entre os subtipos histológicos, o adenocarcinoma e o carcinoma de células escamosas são os mais prevalentes e apresentam diferenças morfológicas sutis, o que torna sua classificação complexa e dependente da interpretação do patologista. **Objetivos:** Este estudo avalia o desempenho das arquiteturas EfficientNet_B0, EfficientNet_V2 e MobileNet_V2 na classificação multiclasse de imagens histológicas pulmonares. Também investiga o impacto de distintas estratégias de aumento de dados, comparando transformações simples e combinações mais complexas, a fim de identificar configurações que maximizem a acurácia e o equilíbrio entre classes. **Método:** O conjunto de dados foi composto por 307 imagens histológicas de tecido pulmonar desbalanceadas entre classes: 90 de adenocarcinomas, 151 de tecidos normais e 66 de carcinomas de células escamosas. As imagens foram submetidas a diferentes estratégias de aumento de dados — flip, rotation, color, crop, combined (integração das quatro anteriores) e combined_v2 (versão com parâmetros ajustados). Cada estratégia foi aplicada separadamente ao conjunto de treinamento, mantendo a validação fixa. Os modelos foram treinados com pesos pré-treinados no ImageNet, em dois modos de fine-tuning (PT-FC e PT-ALL), e avaliados por acurácia e F1-score macro. **Resultados:** A MobileNet_V2 apresentou o melhor desempenho, com 92,52% de acurácia e F1-score de 0,91. A EfficientNet_V2 obteve 87,85% e 0,86, enquanto a EfficientNet_B0 alcançou 86,28% e 0,86, sem aumento. O uso de aumento de dados mostrou-se essencial para melhorar a generalização, especialmente em classes minoritárias. **Conclusão:** A escolha da arquitetura e das estratégias de aumento de dados influencia significativamente o desempenho em classificação multiclasse de imagens histológicas pulmonares. Estratégias combinadas de aumento favoreceram a generalização dos modelos, enquanto a MobileNet_V2 destacou-se pelo equilíbrio entre precisão e eficiência computacional, demonstrando potencial para aplicações em patologia digital auxiliada por aprendizagem profunda.

Trilha: Trabalho de Graduação.

Avaliação de Efeito do Fine-Tuning em Modelos CNNs para a Classificação de Lesões Histológicas

Vitória Silva¹, Luana Borges¹,
Marcelo Nascimento¹, Domingos Oliveira²

¹Universidade Federal de Uberlândia (UFU)

²Instituto Federal de São Paulo (IFSP)

{vitoria.fernandes, luana.borges1, marcelo.nascimento}@ufu.br
domingos.latorre@gmail.com

Introdução: A análise de imagens histológicas com redes neurais convolucionais (CNNs) tem se mostrado oportuna no apoio ao diagnóstico de doenças. O desempenho das CNNs depende da quantidade de dados e do ajuste dos hiperparâmetros. O estágio de fine-tuning surge como uma alternativa eficiente ao treinamento scratch, no qual os pesos da rede são inicializados aleatoriamente, pois permite adaptar modelos pré-treinados em grandes bases de dados (como ImageNet) aos contextos histológicos, nos quais os conjuntos de imagens são menores e mais específicos. Assim, compreender como diferentes estratégias de fine-tuning influenciam o aprendizado é essencial para otimizar o uso dessas arquiteturas em diagnósticos médicos. **Objetivos:** Investigar o efeito do fine-tuning em CNNs aplicadas à classificação de imagens histológicas de displasias da cavidade oral e de tecidos pulmonares, avaliando a parametrização e a capacidade de generalização na identificação dos estágios de lesão. **Método:** Foram investigadas neste trabalho as arquiteturas VGG13, VGG19 e SqueezeNet, sob os modos de treinamento: Scratch, Fine-Tuning apenas do classificador, Fine-Tuning com descongelamento parcial e Fine-Tuning completo. Para cada combinação, variaram-se dois valores de learning rate (0.0001 e 0.001) e as taxas de dropout (0.0 e 0.5), com aumento de dados padronizado. Os experimentos foram conduzidos nos datasets histológicos: um de displasia oral, com duas classes balanceadas, e outro de tecido pulmonar, com três classes desbalanceadas. As métricas avaliadas incluíram acurácia, perda de validação e F1-Score macro. Para cada modo e associação de learning rate e dropout, foram analisados os resultados. **Resultados:** O desempenho das CNNs variou conforme o grau de fine-tuning e os hiperparâmetros aplicados. Em geral, o fine-tuning completo apresentou os melhores resultados em conjuntos balanceados, enquanto o parcial se mostrou mais estável em datasets desbalanceados, reduzindo o overfitting e melhorando a consistência das previsões. As redes VGG mostraram maior sensibilidade às taxas de learning rate e dropout, enquanto a SqueezeNet destacou-se pela eficiência computacional e bom equilíbrio entre desempenho e tempo de treinamento. **Conclusão:** O fine-tuning mostrou ser uma boa estratégia para aprimorar o desempenho de CNNs em imagens histológicas, desde que o grau de ajuste e os hiperparâmetros sejam definidos ajustado para o dataset investigado. Estratégias mais robusta de fine-tuning tendem a favorecer conjuntos menores e balanceados, enquanto abordagens parciais são mais indicadas para cenários desbalanceados. O equilíbrio entre especialização e generalização depende diretamente da arquitetura escolhida, da complexidade tissular e da variabilidade dos dados disponíveis.

Trilha: Trabalho de Graduação.

Avaliação dos Modelos EfficientNetV2 e MobileNetV2 na Classificação de Imagens Histológicas de Displasia da Cavidade Oral

Daniel Gonçalves¹, Gustavo Miranda¹,
Marcelo Nascimento¹, Domingos Oliveira²

¹Universidade Federal de Uberlândia (UFU)

²Instituto Federal de São Paulo (IFSP)

{danielbgoncalves, gustavo.miranda, marcelo.nascimento}@ufu.br
domingos.latorre@gmail.com

Introdução: As displasias são lesões que comprometem o desenvolvimento celular dos tecidos da cavidade oral. Elas apresentam características morfológicas como crescimento anormal, variações de formato e hiperchromasia. Constituem lesões precursoras do câncer bucal, cuja progressão ocorre em aproximadamente 6% a 36% dos casos. Neste estudo, foi investigado o nível de displasia severa em comparação com tecidos saudáveis. A análise tradicional dessas lesões, feita exclusivamente por observação manual, pode ser complexa devido à carga de trabalho e à subjetividade associada à avaliação do especialista. **Objetivos:** Este trabalho propõe um estudo de modelos baseados em redes neurais convolucionais (CNNs), especificamente as arquiteturas EfficientNetV2 e MobileNetV2, para a classificação de imagens histológicas de displasias da cavidade oral. Esses modelos foram escolhidos por apresentarem baixo custo computacional, característica importante para análises em imagens histológicas coradas em H&E. A MobileNetV2 destaca-se por sua leveza e eficiência em conjuntos de dados reduzidos, enquanto a EfficientNetV2 demonstra excelente capacidade de generalização e treinamento otimizado. **Método:** Os modelos das CNNs foram avaliados considerando os parâmetros de aumento de dados, Dropout e taxa de aprendizagem, utilizando um banco de 228 imagens de tecidos saudáveis e displasia. Foram realizadas 200 execuções de treinamento dos modelos. **Resultados:** Os melhores resultados obtidos para as métricas F1-score e acurácia foram ambos iguais a 1, o que indica desempenho perfeito na classificação das imagens de validação. Esses resultados foram alcançados tanto pelo modelo EfficientNetV2, com fine-tuning em todas as camadas e aplicação de técnica de data augmentation, quanto pelo MobileNetV2 sob mesmas configurações experimentais. **Conclusão:** Esses resultados reforçam que as arquiteturas EfficientNet2 e MobileNetV2 apresentam desempenho altamente promissor para o auxílio automatizado na investigação e diagnóstico de lesões displásicas. Portanto, poderiam contribuir para sistemas de suporte à decisão clínica em patologia digital.

Trilha: Trabalho de Graduação.

Avaliando Diferentes Metodologias para o Ensino de Computação Gráfica no Nível Superior

Carla Silva

Universidade Federal de Uberlândia (UFU)

carla.az@ufu.br

Introdução: Na reestruturação da grade curricular do curso de Ciência da Computação pela Faculdade de Ciência da Computação (FACOM) da UFU, Computação Gráfica (CG) foi incluída como disciplina obrigatória. Considerando que essa matéria foi ofertada poucas vezes, como componente optativo, é necessário estruturar a forma que esse curso deve ser lecionado, considerando as tecnologias atuais e as metodologias aplicadas nas demais universidades e nos principais livros e materiais usadas para o aprendizado de Computação Gráfica. **Objetivos:** Essa pesquisa tem como objetivo levantar as diferentes abordagens de ensino de CG, entendendo como essa matéria é comumente abordada, assim como revisar os materiais e plano de ensino usados na última vez que a disciplina foi ofertada. Adicionalmente o trabalho tem o fim de complementar a formação da aluna pesquisadora, com o estudo do conteúdo de Computação Gráfica. **Método:** Para levantar as abordagens mais relevantes para o ensino da matéria, foram buscadas as melhores universidades do país, filtrando apenas as que ofertam atualmente a disciplina de Computação Gráfica. Para isso, foi usado o índice Conceito Preliminar de Curso (CPC), feito e disponibilizado pelo MEC. Com base nas 10 melhores universidades, foram buscadas as suas respectivas ementas de CG e, com base nas bibliografias, foram reunidas as principais referências e materiais utilizados. Com base nesses dados, podemos entender as principais metodologias aplicadas e as tecnologias usadas, sendo possível comparar a atual maneira que o conteúdo é abordado na FACOM com as abordagens usadas no Brasil. **Resultados:** As 10 melhores faculdades, segundo o índice CPC foram a IFG, IFSP, UFABC, UFAL, UFAM, UFERSA, UFJF, UFV, UNICAMP e UNIFESP e, após juntar as bibliografias, foram encontrados 49 materiais distintos, sendo que apenas 11 deles foram referenciados mais de uma vez. Filtrando apenas os materiais que compunham parte da bibliografia obrigatória da ementa, temos 20 materiais utilizados. **Conclusão:** O trabalho, até o momento, foi capaz de identificar os cursos mais relevantes de Ciência da Computação do Brasil, e reuniu as principais bases para o ensino de Computação Gráfica nas universidades brasileiras. Entretanto, ainda é necessária análise das referências levantadas, reunindo as tecnologias que elas abordam, os conteúdos que trazem, e em que ordem, além das metodologias que cada uma adota, para aperfeiçoar a atual ementa disposta pela FACOM e fornecer uma base para que os docentes possam elaborar seus planos de ensino para a disciplina de Computação Gráfica.

Trilha: Trabalho de Graduação.

Calibração Evolutiva de um Modelo de Autômatos Celulares e Otimização no Posicionamento de Sinalização

Pedro Simoni, Luiz Gustavo A. Martins

Universidade Federal de Uberlândia (UFU)

{pedro.simoni.sc, lgamartins}@ufu.br

Introdução: A simulação computacional de pedestres, essencial ao planejamento urbano e segurança, enfrenta desafios na modelagem do comportamento adaptativo e na calibração. Modelos de Autômatos Celulares (AC) tratam agentes com comportamentos estáticos. Metodologias de calibração (ex: Algoritmos Genéticos, AGs) focam em métricas globais (tempo de evacuação), falhando em capturar dinâmicas internas e padrões emergentes (auto-organização de fluxos), gerando baixa fidelidade comportamental. O problema reside na lacuna entre otimização global e replicação de comportamentos coletivos realistas.

Objetivos: O objetivo principal é desenvolver e validar uma metodologia de otimização evolutiva (AG) capaz de calibrar simulações de AC para replicar comportamentos realistas em nível global e local. Objetivos específicos: (1) propor um AG que configure simultaneamente os parâmetros comportamentais do AC (baseados na literatura) e (2) otimize o posicionamento de sinalizações no ambiente, visando o melhor auxílio aos pedestres; (3) validar esta abordagem através de uma função de aptidão multifacetada que meça tanto a eficiência (tempo total) quanto padrões de fluxo locais.

Método: O método foca no desenvolvimento do AG. O simulador AC (em Go) será modelado com base em trabalhos estabelecidos da literatura, servindo como uma base robusta para a otimização. A inovação metodológica está no AG: seu cromossomo codificará duas classes de informação: (1) os parâmetros de comportamento do AC (ex: nível de confusão) e (2) a configuração e posição das sinalizações no grid. A função de aptidão (fitness) avaliará cada indivíduo (configuração de AC + sinais) em dois eixos: eficiência da evacuação (global) e fidelidade da dinâmica interna (padrões locais, como auto-organização e “stop-and-go”).

Resultados: A análise focará na validação da hipótese: calibração por padrões internos e otimização de sinais geram fluxos mais realistas e eficientes. Três cenários (corredor, gargalo, obstáculos) serão usados. Espera-se que o AG identifique configurações ótimas de sinalização que, combinadas aos parâmetros do AC, (1) melhorem métricas globais (tempo) e (2) induzam padrões de fluxo locais desejáveis, como faixas de fluxo e redução de “stop-and-go”. A contribuição será demonstrar que o AG pode otimizar ativamente o ambiente (sinais) para melhorar o comportamento.

Conclusão: A discussão comparará os padrões emergentes com a literatura. A relevância da solução é mover o paradigma da calibração para a “otimização de design de ambiente”, onde o AG atua como uma ferramenta de planejamento. A vantagem é gerar modelos preditivos que não apenas simulam, mas ativamente propõem melhorias (sinalização) para dinâmicas humanas plausíveis. A principal diferença é a tese de que o AG pode otimizar tanto o agente quanto o ambiente. Uma limitação é o foco reduzido em novos comportamentos do AC, priorizando a otimização do ambiente, o que aponta para trabalhos futuros.

Trilha: Trabalho de Graduação.

Capacitando Profissionais de Saúde e Educadores no Diagnóstico de Doenças Raras e na Interpretação Genômica por meio da Bioinformática

Davi Campos, Murilo Beppler, Anderson Santos

Universidade Federal de Uberlândia (UFU)

{davimcampos09, murilobeppler, santosandr}@ufu.br

Introdução: O diagnóstico de doenças raras permanece um desafio significativo para a comunidade científica e médica, não apenas pela complexidade dos dados genômicos, mas também pelo fato de que, em muitos casos, o diagnóstico demora anos para ser concluído — ou sequer acontece. Essa dificuldade decorre da diversidade de manifestações clínicas, da escassez de especialistas e da falta de integração entre dados clínicos e genômicos. Embora os avanços em tecnologias de sequenciamento tenham ampliado o acesso a informações moleculares, a análise e anotação de variantes genéticas ainda exigem conhecimento especializado em bioinformática. Essa limitação técnica dificulta a utilização eficiente dos dados genômicos no processo diagnóstico e restringe a adoção de abordagens baseadas em sequenciamento em ambientes clínicos. **Objetivos:** Desenvolver uma plataforma web integrada a um pipeline bioinformático capaz de auxiliar profissionais de saúde e pesquisadores na anotação e interpretação de dados de exomas relacionados a doenças raras. A proposta visa oferecer uma interface de visualização simples e funcional, permitindo que os resultados das análises genômicas sejam consultados e interpretados de forma mais ágil e acessível. **Método:** O pipeline bioinformático foi estruturado para realizar a anotação de variantes a partir de dados de exoma, integrando etapas de filtragem, priorização e cruzamento com bases públicas de referência, sendo executado via linha de comando. Em paralelo, foi desenvolvido um site que atua como ambiente de visualização e organização dos resultados, permitindo o acesso intuitivo às anotações geradas, bem como às informações clínicas associadas a pacientes e amostras. Diferentemente de ferramentas tradicionais que exigem execução técnica e não oferecem uma camada de apresentação consolidada, a plataforma proposta busca otimizar o acesso aos resultados e favorecer a interpretação por profissionais da área da saúde. **Resultados:** Foi obtido um protótipo funcional que possibilita a navegação e exploração dos resultados de anotações genômicas, com foco em casos simulados de doenças raras. O sistema mostrou-se eficiente para a organização e exibição de informações complexas de forma clara, contribuindo para a interpretação dos resultados e para o acompanhamento integrado de dados genéticos e clínicos. **Conclusão:** O desenvolvimento da plataforma representa um avanço inicial na criação de ferramentas de apoio ao diagnóstico baseadas em bioinformática. Ao combinar um pipeline de anotação validado com uma interface de visualização acessível, o trabalho contribui para tornar a interpretação genômica mais prática e compreensível para profissionais de saúde e pesquisadores, abrindo caminho para futuras validações com dados reais e para o uso clínico em larga escala.

Trilha: Trabalho de Graduação.

Classificação Hierárquica Aplicada a Sistemas de Detecção de Intrusões em Redes IoT

Giovanna Martins, Rodrigo Miani

Universidade Federal de Uberlândia (UFU)

{gimartins, miani}@ufu.br

Introdução: A evolução das tecnologias digitais e a crescente adoção da Internet das Coisas ampliaram a superfície de ataque das redes, aumentando a ocorrência de intrusões e exigindo Sistemas de Detecção de Intrusão mais eficazes. Classificadores planos tradicionais apresentam limitações importantes: modelos binários oferecem alta sensibilidade, mas não distinguem tipos específicos de ataques, enquanto modelos multiclasse sofrem degradação de desempenho conforme o número de categorias cresce. Assim, o problema científico central envolve equilibrar sensibilidade e granularidade na detecção de intrusões em redes IoT.

Objetivos: O trabalho busca desenvolver um classificador hierárquico capaz de combinar a eficiência da detecção binária com a capacidade discriminativa do multiclasse. Além disso, pretende-se analisar diferentes organizações hierárquicas, avaliar o impacto da decomposição do problema e comparar o desempenho preditivo e computacional com abordagens tradicionais, evidenciando ganhos e limitações. **Método:** A metodologia envolveu a extração e o pré-processamento de uma amostra estratificada do conjunto de dados CICIoT2023, a construção de classificadores planos e hierárquicos utilizando Random Forest como modelo base e a avaliação dos classificadores criados. Foram propostas duas hierarquias: BDM (Binário, DDoS e Multiclasse), focada na separação inicial de tráfego benigno e ataques volumétricos, e BRDM (Binário, Recon, DDoS e Multiclasse), que adiciona um nível específico para ataques de reconhecimento. As abordagens foram comparadas quanto ao desempenho preditivo, avaliado por métricas como acurácia, precisão, recall e F1-score, calculadas a partir da matriz de confusão, bem como pelo custo computacional, considerando o tempo de treinamento, o tempo de teste e a quantidade de memória RAM ocupada pelo classificador. **Resultados:** O classificador BDM aumentou a sensibilidade geral sem prejudicar a distinção entre categorias, mas ainda mostrou fragilidades em ataques minoritários. O BRDM aprimorou esses resultados, especialmente em classes desafiadoras, demonstrando maior equilíbrio e interpretabilidade, embora com custo computacional mais elevado. A comparação com métodos relacionados evidenciou que a estratégia hierárquica reduz ambiguidades típicas do multiclasse e aumenta a robustez da detecção. **Conclusão:** A classificação hierárquica mostrou-se uma solução promissora para detecção de intrusões em IoT, oferecendo melhor equilíbrio entre sensibilidade e detalhamento dos ataques. A organização do BRDM destacou-se por mitigar confusões entre categorias semanticamente próximas, contribuindo para diagnósticos mais precisos e úteis para defesa cibernética. Ainda assim, desafios relacionados ao desbalanceamento e ao custo computacional permanecem, indicando oportunidades para otimizações futuras.

Trilha: Trabalho de Graduação.

Computação Inclusiva para Crianças e Jovens Neurodivergentes

Breno Carrijo Pereira, Ariany Silva,

Ana Cláudia Martinez, Giullia de Menezes

Universidade Federal de Uberlândia (UFU)

{brenocarrijo, ariany.silva, anacmartinez, giullia.rodriques}@ufu.br

Introdução: A exclusão digital de jovens neurodivergentes (TEA) frequentemente ocorre devido à “inacessibilidade pedagógica”, já que métodos tradicionais de ensino falham em atender suas necessidades. A Computação Desplugada surge como uma abordagem promissora por ser lúdica ao ensinar os fundamentos do Pensamento Computacional. No contexto do projeto de extensão “Mentes Brilhantes”, este trabalho de Iniciação Científica foca no desenvolvimento, documentação e validação da metodologia e dos materiais didáticos. **Objetivos:** O objetivo geral desta pesquisa é estruturar e validar um conjunto de artefatos pedagógicos e uma metodologia de ensino de Pensamento Computacional Desplugado adaptada às necessidades de crianças e jovens autistas. Os objetivos específicos incluem: Desenvolver um conjunto de materiais didáticos alinhados aos brinquedos pedagógicos e com suporte de terapeutas; Documentar a metodologia de ensino para conceitos de PC (sequência lógica, algoritmos, depuração); e Propor um modelo de avaliação qualitativa para a eficácia da abordagem no engajamento e aprendizado dos participantes. **Método:** Esta pesquisa utilizará uma metodologia de pesquisa-ação, planejada para ocorrer simultaneamente ao projeto de extensão. O processo envolverá o design e a criação iterativa dos materiais didáticos que serão validados por psicólogos e pedagogos parceiros antes da aplicação. As atividades serão estruturadas para aplicação em oficinas semanais, na proporção de um monitor para dois participantes, visando garantir o acompanhamento individualizado. A coleta de dados para a IC será realizada por meio de instrumentos qualitativos, como diários de bordo da equipe, observação participante e formulários de feedback semiestruturados a serem aplicados aos pais e terapeutas acompanhantes após a execução dos ciclos. **Resultados:** Como o projeto está em fase de planejamento, esta seção descreve os resultados esperados. Espera-se entregar, como produtos concretos, as primeiras versões das apostilas visuais e planos de aula (focados em lógica e algoritmos) validadas por terapeutas. Almeja-se que a aplicação futura demonstre alto engajamento e que os materiais visuais, focados na previsibilidade, se mostrem essenciais para a confiança do público TEA, assim como a metodologia de cocriação com terapeutas será fundamental para o ajuste às necessidades sensoriais do grupo. **Conclusão:** Este trabalho buscará evidenciar o potencial da Computação Desplugada como uma poderosa ferramenta de inclusão sociodigital. A estruturação de uma metodologia e o desenvolvimento de materiais didáticos específicos serão as contribuições diretas desta IC, oferecendo um modelo replicável para outras instituições. Espera-se demonstrar que o profissional de Sistemas de Informação pode atuar para além do desenvolvimento de software, contribuindo ativamente no desenvolvimento de soluções pedagógicas que promovem a redução das desigualdades.

Trilha: Trabalho de Graduação.

DataRU: Ciência de dados aplicada aos Restaurantes Universitários - Resultados Finais

Matheus G. S. Resende, Fabíola S. F. Pereira

Universidade Federal de Uberlândia (UFU)
{matheus-resende, fabiola.pereira}@ufu.br

Introdução: A gestão dos Restaurantes Universitários (RUs) da UFU possui uma pesquisa de satisfação anual, com baixa adesão (cerca de 2%) e sem monitoramento diário. A proposta inicial desta pesquisa era solucionar isso com a plataforma “Data RU”, para coleta de avaliações diárias dos estudantes, de 1 a 5 estrelas. Contudo, barreiras institucionais, técnicas e burocráticas inviabilizaram a implementação do app. Assim, o projeto foi redefinido em alinhamento com a gestão, pivotando da coleta de novos dados para a análise de dados existentes, que estavam inacessíveis em arquivos PDF. **Objetivos:** O objetivo original era implementar o app “Data RU” para monitoramento diário, visando criar rankings de pratos e prever filas. Após a mudança, os objetivos foram: 1) Impulsionar uma comunidade com avaliação das refeições pelo Whatsapp. 2) Estruturar os dados de consumo e cardápios de arquivos PDF em um banco de dados funcional. 3) Aplicar ciência de dados exploratória nessa base para criar scripts SQL capazes de gerar painéis para a gestão (PROAE/DIVRU). 4) Demonstrar o valor prático da análise de dados como ferramenta de apoio à decisão, com a gestão da Reitoria. **Método:** A metodologia planejada era a coleta primária de dados via avaliações no app. A metodologia executada foi um processo clássico de ciência de dados focado intensamente em engenharia de dados. O maior desafio foi a obtenção e fusão dos dados. Foi realizado um “data scrapping” (raspagem de dados) manual e semiautomático para extrair todas as informações de consumo e cardápios dos PDFs. Na sequência, os dados foram modelados e padronizados em um banco de dados. Por fim, foi desenvolvido scripts SQL para realizar as análises, cujos resultados foram apresentados em reuniões com a gestão superior. **Resultados:** Os principais produtos técnicos gerados foram uma base de dados estruturada do consumo e cardápios dos RUs e um conjunto de scripts SQL para análise de dados. As análises possibilitadas por estes produtos geraram insights valiosos, como rankings de pratos e médias de consumo, que foram formalmente apresentados à administração dos RUs. O projeto estabeleceu um canal de diálogo efetivo com a PROAE e a DIVRU. A pesquisa obteve reconhecimento acadêmico e nacional, sendo apresentado em eventos científicos e recebendo premiação na Mostra Científica da 14ª Bienal da UNE. **Conclusão:** Esta pesquisa cumpriu seu objetivo de aplicar a ciência de dados para gerar valor real à comunidade acadêmica, adaptando seu escopo original. O projeto superou barreiras burocráticas, provando o valor da engenharia de dados ao transformar arquivos estáticos em uma base de dados rica e analisável. O reconhecimento em eventos científicos validou a relevância social do tema e técnica da pesquisa. Obtive conhecimentos práticos nos fundamentos da Ciência de Dados, trabalhando com proatividade junto à minha orientadora. Como trabalho futuro, um artigo científico será preparado para submissão.

Trilha: Trabalho de Graduação.

Deep Learning em Patologia Digital: Avaliação de ResNet e MobileNet na Análise Histológica de Displasia Oral e Pulmão

Davi Soares¹, Luis Paim¹, Domingos Oliveira², Marcelo Nascimento¹

¹Universidade Federal de Uberlândia (UFU)

²Instituto Federal de São Paulo (IFSP)

{davi.tiago, marcelo.nascimento}@ufu.br

{luisfeliipepaimti, domingos.latorre}@gmail.com

Introdução: A displasia epitelial oral é um precursor comum do carcinoma de células escamosas da cavidade bucal, cuja progressão para malignidade varia entre 6% e 36%. Apesar dos avanços terapêuticos, esse tipo de tumor permanece entre as neoplasias mais incidentes e com baixa taxa de sobrevida a longo prazo. As displasias apresentam alterações morfológicas progressivas do epitélio, evoluindo de hiperplasia inicial para estágios leves, moderados e severos. De modo análogo, o carcinoma pulmonar é uma neoplasia agressiva associada a elevada mortalidade e desafios diagnósticos. O exame histopatológico é essencial para o diagnóstico e conduta clínica, mas a classificação dessas lesões é complexa e depende da experiência do patologista, o que pode gerar subjetividade. **Objetivos:** Desenvolver e avaliar uma metodologia de classificação automatizada de imagens histopatológicas de displasia epitelial oral e carcinoma pulmonar utilizando redes neurais convolucionais (CNNs), a fim de auxiliar especialistas na análise e reduzir a subjetividade diagnóstica. **Método:** Foram investigadas as arquiteturas ResNet18, ResNet34 e MobileNetV3 Small, selecionadas por sua eficácia em tarefas de classificação visual e pela capacidade de equilibrar desempenho e eficiência computacional, características cruciais em sistemas de apoio diagnóstico. A metodologia baseou-se em transfer learning, utilizando pesos pré-treinados para acelerar a convergência e melhorar a generalização dos modelos. As camadas finais foram adaptadas para a classificação das duas patologias, considerando imagens histológicas obtidas de bases específicas para cada tipo de lesão. **Resultados:** Os modelos ResNet18 e ResNet34 alcançaram acurácias de 98,55% na classificação da displasia oral e 93,55% no carcinoma pulmonar. A MobileNetV3 Small, mesmo com arquitetura mais leve, apresentou resultados competitivos, com 97,10% e 88,17%, respectivamente. Esses resultados demonstram o potencial das abordagens baseadas em Deep Learning na análise histológica, destacando o equilíbrio entre precisão e custo computacional obtido pelas arquiteturas avaliadas. **Conclusão:** As redes ResNet e MobileNetV3 mostraram desempenho consistente na classificação de displasia oral e carcinoma pulmonar, evidenciando o potencial do Deep Learning como ferramenta de suporte ao diagnóstico patológico. A adoção dessas técnicas pode contribuir para maior padronização e eficiência na análise de lâminas histológicas.

Trilha: Trabalho de Graduação.

Dependabilidade do Sistema Operacional na Otimização de Modelos de Inteligência Artificial: Um Estudo de Caso Aplicado à Cibersegurança

Maria Eduarda Teixeira, Diego Molinos

Universidade Federal de Uberlândia (UFU)
{maria.teixeira27, diego.molinos}@ufu.br

Introdução: A crescente adoção de modelos de inteligência artificial (IA) no campo da cibersegurança tem proporcionado avanços significativos na detecção de ameaças cibernéticas, em especial ransomwares. No entanto, a implementação eficiente desses modelos em ambientes reais ainda enfrenta desafios relevantes relacionados ao desempenho computacional, principalmente em ambientes com restrições de recursos computacionais tais como soluções embarcadas, dispositivos IoT e até mesmo firewalls. Neste contexto, torna-se fundamental considerar a dependência de sistemas operacionais que sustentam a execução desses modelos, uma vez que a orquestração entre o sistema operacional e o framework de aprendizado de máquina influencia diretamente a estabilidade, a disponibilidade e a previsibilidade do desempenho das soluções baseadas em IA. **Objetivos:** Diante desse cenário, o objetivo deste trabalho é considerar fatores como análise de processos/threads e alocação de recursos de memória computacional durante a execução dos modelos, para assim analisar o consumo de recursos e a previsibilidade de execução, otimizando seu desempenho mesmo em plataformas de baixo custo ou com processamento limitado. **Método:** Metodologicamente, esta pesquisa é de caráter experimental e exploratório e, encontra-se estruturada em cinco etapas principais, são elas: (a) estudo do estado da arte sobre detecção de ransomwares e machine learning; (b) escolha e pré-processamento do dataset ((CIC-MalMem-2022)); (c) treinamento de classificadores do tipo Árvore de Decisão, Naive Bayes e Rede Neural; (d) exportação e teste dos modelos respondendo as inferências em um ambiente controlado com Linux, inspecionando o desempenho computacional, como abertura e fechamento de threads e alocação de memória; e (e) análise comparativa dos resultados obtidos, correlacionando métricas de desempenho computacional (tempo de inferência, uso de memória RAM e CPU) com as taxas de acurácia dos classificadores, visando identificar o modelo mais adequado para implementação em ambientes com recursos limitados. **Resultados:** Até o momento a etapa (a), (b) e (c) já foram realizadas, os resultados preliminares apresentam um desempenho significativamente superior dos modelos de Naive Bayes e Rede Neural em relação à Árvore de Decisão, com acurácias de 99,26% e 100%, respectivamente, frente a apenas 48,75% da árvore. **Conclusão:** Esses resultados indicam que os modelos mais complexos foram capazes de generalizar melhor os padrões entre amostras benignas e de ransomware, espera-se que com a execução das etapas (d) e (e) seja possível observar o real custo operacional da execução dos modelos no que tange à abertura de processos/threads e memória do sistema operacional, permitindo assim avaliar a viabilidade prática de cada abordagem em ambientes de produção e o equilíbrio entre desempenho preditivo e custo computacional.

Trilha: Trabalho de Graduação.

Desafios no Pré-processamento de Dados Clínicos para Aplicação em Inteligência Artificial: Um Estudo de Caso em Reumatologia

César Cardoso, Fernanda Santos,
Alessandra Paulino, Gabriel Teles, Daniel Caetano

Universidade Federal de Uberlândia (UFU)
{cesar.cardoso, fmcsantos}@ufu.br
{alessandra, gabriel.teles1, daniel.caetano}@ufu.br

Introdução: A crescente demanda por soluções tecnológicas na saúde tem impulsionado o desenvolvimento de ferramentas que aprimorem diagnósticos e decisões clínicas. Contudo, a precisão diagnóstica pode ser comprometida por limitações na experiência dos profissionais e pela falta de clareza nas informações dos pacientes. A literatura destaca a dificuldade de estruturar adequadamente dados clínicos, o que prejudica a criação de corpos aplicáveis a algoritmos de Inteligência Artificial (IA) voltados ao suporte à decisão médica. **Objetivos:** Este trabalho vem reforçar os relatos sobre os desafios na estruturação de dados clínicos com base na prova de conceito que está em desenvolvimento num módulo de apoio à decisão terapêutica de pacientes atendidos na Reumatologia da UFU, já aprovado pelo Comitê de Ética sob o número CAAE: 78174224.8.0000.5152. O estudo utiliza dados estruturados cedidos pelo Departamento de Reumatologia do Hospital de Clínicas da Universidade Federal de Uberlândia (HC-UFU), anonimizados. **Método:** Este estudo compreende as duas primeiras etapas do projeto que compreende: (1) Coleta e seleção de dados; (2) Pré-processamento ; (3) Implementação de algoritmos de PLN; e (4) Análise exploratória dos dados (EDA). A primeira etapa foi realizada a inclusão de prontuários de pacientes maiores de 18 anos e a exclusão de casos em tratamento ativo. A segunda etapa realizou a limpeza, padronização e tratamento de ruídos e de dados ausentes do corpus em construção. **Resultados:** Os dados analisados apresentaram alta inconsistência e baixa correlação, exigindo intenso trabalho de pré-processamento. Foram desenvolvidos algoritmos em Python, para leitura e análise dos arquivos que continham informações sobre anamnese, receitas, evoluções e exames. Após examinar mais de 40 mil registros, apenas um conjunto de dados válido pôde ser consolidado, limitando a aplicação de algoritmos de IA, que dependem de grandes volumes de informação. Essa limitação evidencia o impacto da qualidade dos dados na construção de sistemas de apoio à decisão clínica. **Conclusão:** O projeto enfrenta desafios significativos decorrentes da desorganização e inconsistência dos dados, situação comum na área da saúde. Apesar dos esforços de processamento e consolidação, os resultados não permitiram o avanço para as etapas seguintes. Logo, este estudo reforça a necessidade de bases de dados mais estruturadas e padronizadas para viabilizar o desenvolvimento de soluções computacionais eficazes que contribuam para auxílios aos profissionais da saúde.

Trilha: Trabalho de Graduação.

Detecção de Diabetes Mellitus a partir da Espectroscopia de Infravermelho com Transformada de Fourier via Aprendizado de Máquina

Kamily Cristina Gomes, Douglas C. Caixeta,
Murillo G. Carneiro, Robinson S. Silva

Universidade Federal de Uberlândia (UFU)

{kamily.gomes, mgcarneiro}@ufu.br, {caixetadoug, robinsonsabino}@gmail.com

Introdução: Métodos de Machine Learning (ML) estão cada vez mais sendo aplicados para aprender padrões existentes em grandes conjuntos de dados e prever novos com uma maior exatidão. A identificação precoce da Diabetes Mellitus (DM) é essencial, pois assim maiores serão as chances de que o paciente tenha um bom tratamento e qualidade de vida a posteriori. Diante disso, o uso da técnica de espectroscopia de infravermelho com transformada de Fourier por reflectância total atenuada (ATR-FTIR), que gera o espectro de infravermelho da amostra analisada, tem sido de extrema importância para a geração dos espectros de forma menos invasiva, uma vez que, neste trabalho utilizamos como biofluido para a extração de informações a saliva. **Objetivos:** O objetivo da pesquisa consiste em analisar métodos convencionais de ML, tais como: Discriminant Analysis (LDA), C-Support Vector Classification (C-SVC), Random Forest (RF), Logistic Regression (LR) e uma arquitetura de rede neural convolucional (CNN) a fim de classificar mais precisamente os espectros ATR-FTIRs. **Método:** A base de dados utilizada consiste em 215 espectros, dos quais 169 são de indivíduos portadores de DM e os outros 46 pertencentes ao grupo de controle. Cada espectro possui 3596 atributos, os quais foram utilizados pelos modelos para prever sua devida classificação. Esses dados passaram por um pré-processamento visando eliminar ruídos das amostras por meio da correção da linha de base via rubberband e normalização pelo pico da Amida-I, que ajusta todas as intensidades de onda para valores entre 0 e 1. Posteriormente, foi observado um desbalanceamento das classes (DM e controle) e, para tratar isso, técnicas de oversampling e undersampling foram aplicadas e comparadas com o aprendizado a partir das classes desbalanceadas. O procedimento de oversampling consiste em gerar amostras sintéticas para o treino, e o de undersampling em retirar amostras da classe majoritária de forma que as classes fiquem com a mesma quantidade de elementos. Todos os modelos foram treinados e testados com a técnica de validação cruzada estratificada k-fold para avaliar a habilidade do modelo em lidar com dados novos. **Resultados:** Foram analisadas medidas de desempenho tais como: sensibilidade, especificidade e a acurácia balanceada (média entre sensibilidade e especificidade). As melhores métricas dentre os modelos tradicionais de ML foram alcançadas pela LR combinada com oversampling tendo como acurácia balanceada 84,6%. No entanto, o modelo de CNN superou este resultado alcançando 89,1% de acurácia balanceada, sensibilidade de 86,96% e especificidade de 91,3%. **Conclusão:** Portanto, os modelos de ML se mostram promissores para um diagnóstico não invasivo de DM a partir da saliva e seus biomarcadores presentes. A CNN, quando combinada com o procedimento de undersampling, superou todos os outros métodos na métrica de especificidade, demonstrando melhor capacidade de detecção dos verdadeiros negativos para o teste.

Trilha: Trabalho de Graduação.

EcoTech UFU: Tecnologia e Sustentabilidade na Destinação do resíduo eletrônico

Tatiana Tamura, Diego Molinos, Fernando Santil

Universidade Federal de Uberlândia (UFU)

{tatiana.tamura, diego.molinos, fernando.santil}@ufu.br

Introdução: O volume crescente de dispositivos eletrônicos descartados tem ampliado significativamente a geração de resíduos eletrônicos, um fenômeno que representa um dos principais desafios ambientais da atualidade. Substâncias tóxicas presentes em equipamentos descartados de forma inadequada, como chumbo, mercúrio e cádmio, constituem risco direto à saúde humana e à preservação dos ecossistemas. Diante desse cenário, iniciativas de conscientização e de gestão responsável de resíduos-lixo tornam-se essenciais para mitigar impactos ambientais e promover práticas socialmente sustentáveis. No município de Monte Carmelo, observa-se a ausência de um ponto fixo de coleta de resíduos eletrônicos, lacuna que reforça a necessidade de ações institucionais voltadas ao descarte adequado e à educação ambiental. **Objetivos:** O projeto EcoTech UFU tem como propósito promover a conscientização sobre o descarte adequado de resíduos eletrônicos, estimular práticas sustentáveis e consolidar a UFU como referência regional na coleta desses materiais. **Método:** A proposta estrutura-se da seguinte forma: (a) planejamento, reuniões com os parceiros institucionais, cronograma e organização das responsabilidades da equipe, (b) execução, implantam-se pontos fixos de coleta no campus da UFU em Monte Carmelo, promovem-se ações de divulgação e realizam-se atividades educativas voltadas à conscientização sobre o descarte adequado de resíduos eletrônicos, (c) período de coleta, caracterização, quantificação, registro e classificação dos resíduos, garantindo a destinação adequada em parceria com a CODEL e, (d) aplicação de questionários, permitindo analisar o impacto das ações educativas, o engajamento da comunidade e a correlação entre capacitação e volume de resíduos coletados, gerando dados que fundamentam a eficácia científica e social do projeto. **Resultados:** Trata-se de uma pesquisa em andamento. Entre os resultados ambientais, prevê-se a quantificação do volume total de resíduos coletados, sua distribuição por categorias e a estimativa da redução de impactos ambientais decorrente do descarte adequado. No campo educacional, esperam-se ganhos de conhecimento dos participantes por meio de instrumentos pré- e pós-teste, além de mudanças declaradas nas percepções sobre o descarte responsável. O projeto também pretende gerar evidências de associação entre a participação nas atividades educativas e o aumento do descarte correto, permitindo análises estatísticas de correlação. **Conclusão:** O projeto EcoTech UFU demonstra elevada relevância socioambiental ao suprir a falta de pontos de coleta de resíduos eletrônicos em Monte Carmelo e ao promover ações educativas que estimulam práticas de descarte responsável. Espera-se que as intervenções realizadas resultem na redução de impactos ambientais e na conscientização da comunidade. Além disso, os dados produzidos ao longo do projeto poderão subsidiar análises futuras e fortalecer iniciativas de sustentabilidade na UFU e na região.

Trilha: Trabalho de Graduação.

Estudo da MobileNetV2 para Classificação de Câncer de Mama

Gabriel Vitor S. Costa, Ana Cláudia Martinez, Thiago Ribeiro

Universidade Federal de Uberlândia (UFU)

{gabriel.costa1, anacmartinez, tpribeiro}@ufu.br

Introdução: O câncer de mama é uma das principais causas de mortalidade entre mulheres, e o diagnóstico precoce é essencial para aumentar as chances de cura. Com o avanço tecnológico, redes neurais convolucionais (CNNs) vêm se destacando na detecção automatizada de doenças por meio de imagens médicas. Entre elas, a MobileNetV2 se destaca por sua arquitetura leve e eficiente, amplamente utilizada em dispositivos de baixo custo graças às convoluções separáveis em profundidade e blocos residuais invertidos. Estudos recentes demonstram que a MobileNetV2 pode atingir alta acurácia na classificação de imagens de mamografia, sendo uma alternativa promissora para apoiar o diagnóstico clínico.

Objetivos: Avaliar o desempenho da MobileNetV2 na classificação de imagens de câncer de mama, analisando sua viabilidade em contextos clínicos e de pesquisa. Busca-se compreender como o uso de aprendizado por transferência, ajustes de hiperparâmetros e técnicas de pré-processamento influenciam a acurácia do modelo. **Método:** O modelo foi desenvolvido em Python com TensorFlow e Keras, sendo treinado no Google Colab. A MobileNetV2 foi utilizada como extratora de características, empregando pesos pré-treinados no ImageNet, com camadas densas adicionais para a etapa de classificação. Também foram analisadas abordagens complementares, como o uso de wavelets para extração de texturas e o aprendizado federado, que possibilita treinar modelos em bases distribuídas preservando a privacidade dos dados. Os experimentos utilizaram imagens de mamografia de bases públicas, como o MINI-DDSM. **Resultados:** Estudos revisados indicam que a MobileNetV2 pode alcançar acurácias entre 83% e 98 % em diferentes conjuntos de mamografias, variando conforme a qualidade das imagens e o pré-processamento. No experimento realizado, observou-se forte overfitting a partir da terceira época: enquanto a acurácia de treinamento crescia, a de validação caiu acentuadamente, com instabilidade e picos de perda próximos de 9,5. Esse comportamento sugere falha na generalização para dados não vistos, possivelmente devido a uma taxa de aprendizado alta e à ausência de data augmentation. Os próximos ajustes incluirão essas estratégias para obter um modelo mais estável e com melhor capacidade de generalização. **Conclusão:** A MobileNetV2 mostrou-se uma alternativa viável para a detecção de câncer de mama em imagens de mamografia, unindo eficiência computacional e precisão diagnóstica. Sua aplicação é especialmente relevante em sistemas portáteis e soluções de edge computing, favorecendo o uso em ambientes clínicos com recursos limitados. A continuidade deste trabalho envolve expandir a base de dados, testar variações de pré-processamento e investigar modelos híbridos que combinem extração de textura e aprendizado distribuído para aprimorar o desempenho diagnóstico.

Trilha: Trabalho de Graduação.

Estudo de Desempenho de Modelos de Machine Learning na Identificação de Ataques do tipo Port Scanning

Augusto Costa, Diego Nunes

Universidade Federal de Uberlândia (UFU)
{augusto.idalgo, diego.molinos}@ufu.br

Introdução: A segurança da informação é um pilar fundamental na era digital, em que a conectividade onipresente expõe sistemas e dados a um número crescente de ameaças. Diariamente, servidores e aplicações são alvos de inúmeras requisições de rede, muitas delas maliciosas. Métodos tradicionais de detecção, como sistemas baseados em assinaturas, frequentemente se mostram insuficientes para combater ataques zero-day, criando uma necessidade urgente de soluções mais dinâmicas e inteligentes. Entre as ameaças de rede, o Port Scanning merece atenção especial por ser frequentemente utilizado na fase de reconhecimento de ciberataques. Essa técnica visa mapear portas e serviços ativos, revelando potenciais vulnerabilidades que podem ser exploradas posteriormente. Assim, a detecção precisa e antecipada de varreduras de portas é fundamental para prevenir ataques mais complexos e fortalecer a postura de segurança das redes. Nesse cenário, os Sistemas de Detecção de Intrusão (IDS) baseados em Aprendizado de Máquina (Machine Learning) emergem como uma alternativa promissora. **Objetivos:** Este trabalho objetiva avaliar o desempenho de diferentes modelos na detecção de tráfego de rede associado a atividades de varredura de portas (Port Scanning) e também adicionar transparência ao processo de detecção por meio de inteligência artificial explicável (XAI). **Método:** Para treinar os modelos, foi selecionado o dataset público UNSW-NB15, um conjunto de dados abrangente e moderno que contém tráfego de rede benigno e uma variedade de ataques. A fase atual do projeto consiste no tratamento e pré-processamento deste dataset, que envolve a análise dos dados, a normalização e a engenharia de atributos (feature engineering) para selecionar as características mais relevantes do tráfego. Após o término da preparação dos dados, pretende-se implementar e treinar, no mínimo, quatro modelos de machine learning distintos baseados nos algoritmos: Árvores de Decisão, Random Forest, SVM e Redes Neurais. Um pilar central da pesquisa será a avaliação comparativa do desempenho desses modelos e a etapa de racionalização dos modelos por meio de XAI. **Resultados:** até o momento, os resultados preliminares incluem o pré-processamento do dataset, ou seja, balanceamento, normalização, análise de features e geração dos conjuntos de treinamento e teste. **Conclusão:** O diferencial deste projeto não reside apenas na aplicação de ML, mas também na análise comparativa para identificar as abordagens mais eficazes em termos de desempenho. Espera-se que os resultados contribuam para o desenvolvimento de soluções automatizadas capazes de fortalecer as defesas cibernéticas. Academicamente, o projeto contribui para o campo da IA aplicada à segurança, e, na prática, visa oferecer um caminho para reduzir o tempo de detecção de ameaças e minimizar a carga de trabalho de analistas de segurança.

Trilha: Trabalho de Graduação.

Estudo de técnicas para auxiliar no diagnóstico de Câncer de Mama

Pedro Furlan, Thiago Ribeiro

Universidade Federal de Uberlândia (UFU)

{pedro.furlan, tpribeiro}@ufu.br

Introdução: O câncer é uma das causas de óbito mais comuns do mundo, sendo o câncer de mama a neoplasia mais comum entre as mulheres, além de ser o maior causador de mortes nesse grupo. Um dos métodos mais comum para diagnóstico desse tipo de tumor é a mamografia por raio-X. Com o avanço tecnológico, a qualidade da imagem digital melhorou, permitindo que esses exames sejam utilizados para treinar redes neurais e identificar tumores. Essas ferramentas podem ser utilizadas para auxiliar os médicos na identificação desse câncer, como uma segunda opinião, mas sem substituí-los. Isto poderia auxiliar o médico a diagnosticar o tumor em estágios iniciais, que pode ser difícil de identificar a olho nu, porém as redes neurais podem ajudar a confirmar o diagnóstico do profissional quando estiver com dúvida. **Objetivos:** Este trabalho tem como objetivo investigar técnicas de processamento de imagens, aprendizado de máquina e redes neurais para auxiliar no diagnóstico de câncer de mama. Para isso, pretende-se realizar uma pesquisa em artigos e bases de imagens, estudar a aplicação das técnicas em redes neurais, testar os métodos desenvolvidos e analisar os resultados de desempenho. **Método:** Neste projeto serão analisadas imagens de mamografias feitas por raio-X presentes em bancos de imagens públicos. Caso tenha uma quantidade insuficiente de imagens em alguma dessas bases, será feito um aumento de dados para que alcance uma quantidade maior de imagens. Além das imagens serão estudadas técnicas voltadas para o câncer de mama, com o objetivo de melhorar a precisão de redes neurais. As técnicas serão testadas nos bancos de imagens, utilizando métricas para avaliar o resultado das técnicas e analisar o número total de parâmetros da rede e das camadas convolucionais. **Resultados:** Espera-se com esta pesquisa identificar técnicas de processamento de imagens que ajudem a melhorar o desempenho de redes neurais voltadas para identificação de câncer de mama. É esperado que essas técnicas aumentem a precisão dessas redes e diminuam a quantidade de falsos-negativos nesses exames, facilitando e melhorando o diagnóstico de câncer de mama. **Conclusão:** Esse projeto visa analisar e testar técnicas que ajudem a melhorar a precisão de redes neurais para identificar câncer de mama, propondo métodos para auxiliar os médicos a identificar tumores com maior precisão e no estágio inicial. Dessa forma, espera-se poder auxiliar na redução do número de óbitos por essa neoplasia.

Trilha: Trabalho de Graduação.

Evolução de um Sistema para Análise de Planos de Ensino baseado em Inteligência Artificial Generativa

William da Silva, Victor Sobreira

Universidade Federal de Uberlândia (UFU)

{william4ndrade, victor}@ufu.br

Introdução: A verificação de conformidade de planos de ensino quanto às diretrizes institucionais é atividade crítica para a garantia da qualidade, porém é demorada, suscetível a interpretações divergentes e influenciada por formatação heterogênea. Este trabalho dá continuidade e evolui um sistema previamente desenvolvido na UFU, buscando aprimorar a automação e agilidade das análises, ampliar seu suporte a diferentes contextos e formatos de entrada, além de reduzir o tempo dispendido nas análises e as chances de erro e viés, decorrentes de subjetividade e falhas humanas. Esta proposta busca apoiar a análise de conformidade subsidiando as decisões com base em evidências textuais e destacando justificativas legíveis para o revisor humano. **Objetivos:** Definir e validar um método reproduzível para checar, por seção, a aderência de um plano às normas e diretrizes institucionais, gerando relatórios explicáveis. Pretende-se estruturar regras e documentos em um esquema comparável, incorporar recuperação de trechos normativos antes da inferência, oferecer API e interface web para uso do sistema por não especialistas, permitir a troca de modelos de linguagem sem impacto no restante do sistema e registrar evidências e logs para auditoria. **Método:** O método proposto organiza o fluxo em três fases. Na ingestão, realiza extração de texto de PDFs, limpeza e segmentação por seções. Na orquestração, indexa regras e componentes do plano em uma representação comum para permitir alinhamento por item de verificação. Na análise, aplica geração aumentada por recuperação para apresentar ao modelo trechos normativos e trechos do plano, solicitar decisão de conformidade e coletar justificativas. Estão previstos testes automatizados, verificação de ausência de evidências, serviços independentes e endpoints REST para integração com sistemas acadêmicos. A avaliação considerará concordância com revisores humanos, tempo de revisão e estabilidade de respostas, com estudos de ablação para isolar o efeito da recuperação e da engenharia de prompts. **Resultados:** Até o momento, foram conduzidas pesquisas de literatura sobre análise automática de conformidade documental e uso de modelos de linguagem com recuperação, bem como a reexecução do artefato anterior para estabelecimento de um baseline. A análise do trabalho anterior indica bom desempenho em partes estruturadas do plano e maior variabilidade em campos abertos, além de limitações como sensibilidade a ruídos de extração e layout, dependência de regras pouco padronizadas, pouca rastreabilidade das decisões e conjunto reduzido de casos avaliados. Esses achados orientam as escolhas de arquitetura e o desenho do protocolo de avaliação desta evolução. **Conclusão:** Com a evolução deste trabalho, espera-se a produção de um protótipo funcional, validado com relatórios explicáveis, redução do tempo de revisão, maior concordância com avaliadores humanos em itens padronizados, diminuição de decisões sem evidência e maior reprodutibilidade.

Trilha: Trabalho de Graduação.

GenPPI: Análise Comparativa de Interações Proteína-Proteína via Aprendizado de Máquina em Quatro Cepas de *Kosakonia cowanii* isoladas no Triângulo Mineiro

Anderson Santos, Gustavo L. Fonseca, Natanael Avila, Davi Campos

Universidade Federal de Uberlândia (UFU)

{santosandr, gustavoleal, natanael.avila, davimcampos09}@ufu.br

Introdução: As quatro cepas da bactéria *Kosakonia cowanii* recém isoladas de cafeeiros do Triângulo Mineiro (UFU H1/CP160410, UFU D1/CP162577, UFU G119/CP160412 e UFU I87/CP162576) apresentam dupla funcionalidade: agem como Promotoras de Crescimento Vegetal (PGPB) e, ao mesmo tempo, podem causar a queima foliar no café. Diante disso, questiona-se: será possível analisar e comparar as redes de interação proteína-proteína (PPI) dessas quatro cepas para identificar as características moleculares que explicam sua função dual? A ausência de dados de PPI para esta bactéria é um obstáculo para o melhor entendimento de seu funcionamento. Sob essa perspectiva, espera-se que a análise comparativa das redes de PPI, geradas por um modelo de aprendizado de máquina e validadas contra o modelo genérico do software GenPPI, revele diferenças topológicas significativas, apontando para a causa molecular dessa versatilidade funcional. **Objetivos:** O objetivo principal da pesquisa é criar e validar um modelo de aprendizado de máquina específico para o gênero *Kosakonia* no software GenPPI. Subsequentemente, este modelo será utilizado para prever e analisar as redes de PPI das quatro cepas de interesse, a fim de identificar proteínas e módulos funcionais essenciais relacionados à função dual da bactéria. Ademais, busca-se gerar conhecimento inédito no campo da bioinformática sobre as redes de interação de *K. cowanii*. **Método:** Inicialmente, serão coletados genomas da família Enterobacteriaceae de bancos de dados GenBank para treinar um modelo de aprendizado de máquina específico para o gênero *Kosakonia* dentro da plataforma GenPPI. Este modelo, que utilizará diversas características proteicas para alimentar um classificador, será avaliada por meio de validação cruzada. Feito isso, o modelo será usado para prever as interações proteicas nos genomas das quatro cepas de interesse. As interações de alta confiança, selecionadas por um limiar de score, serão usadas para construir as redes, que serão analisadas com os softwares Gephi (visualização) e R/iGraph (análise quantitativa). A análise topológica será integrada a dados de anotação funcional e de potencial de virulência, previamente gerados pelas ferramentas PANNOTATOR e Medpipe. As análises individuais das redes de cada cepa serão comparadas para identificar padrões comuns e divergentes. **Resultados:** Espera-se que, ao final da pesquisa, seja possível identificar com confiabilidade as proteínas, as interações e os módulos funcionais que são a causa da dupla funcionalidade das bactérias analisadas. Acredita-se que este estudo poderá, futuramente, viabilizar o desenvolvimento de bioinsumos mais eficazes e estratégias de controle da doença no café, com base no conhecimento gerado. **Conclusão:** A pesquisa, embora em estágio inicial, demonstra potencial para aprofundar o conhecimento dos mecanismos biológicos que regem bactérias pouco estudadas, combinando de forma inovadora as áreas da computação, biologia e estatística.

Trilha: Trabalho de Graduação.

Geração Automática de Grades Horárias com Restrições de Disponibilidade para um Sistema de Distribuição de Disciplinas

Gustavo Amorim, Victor Sobreira

Universidade Federal de Uberlândia (UFU)
{gustavo.mascarenhas, victor}@ufu.br

Introdução: A alocação de disciplinas em instituições de ensino superior é um processo logístico complexo que impacta diretamente a qualidade acadêmica e a eficiência administrativa. Sistemas acadêmicos anteriores automatizaram partes desse fluxo, carecendo da incorporação dinâmica das restrições de horário docente, exigindo intervenção manual extensiva. Este trabalho propõe a evolução desse sistema, focada no desenvolvimento de um módulo com interface gráfica para geração e otimização da grade de horários, considerando restrições de disponibilidade e preferências do corpo docente. **Objetivos:** O objetivo geral é o desenvolvimento e integração de um módulo ao sistema legado que permita o gerenciamento de preferências e restrições de horário dos professores, automatizando a geração da grade. Os objetivos específicos são: (1) analisar a arquitetura do sistema legado; (2) modelar e implementar a nova estrutura de dados para as restrições docentes; (3) desenvolver uma interface gráfica (GUI) para que professores e coordenadores gerenciem essas restrições; e (4) implementar um algoritmo de alocação que considere as novas restrições. A justificativa reside principalmente na necessidade de otimização do tempo para definição das grades pela coordenação, em especial, na resolução de conflitos. **Método:** A metodologia adotada é a de Pesquisa e Desenvolvimento (P&D). A primeira etapa, já concluída, consistiu na revisão da literatura sobre Problemas de Alocação de Horários (Timetabling Problems) e na escolha das ferramentas tecnológicas. O levantamento de requisitos funcionais e não funcionais foi realizado com base nas limitações do sistema antigo. O desenvolvimento utiliza uma pilha tecnológica composta por React.js para a interface gráfica, Node.js para a API de back-end, e a biblioteca Google OR-Tools como núcleo do motor de otimização para resolver as restrições e preferências. A validação do protótipo será conduzida através de testes unitários, de integração e de usabilidade. **Resultados:** Como resultado, propõe-se o desenvolvimento de um protótipo de software funcional com o novo módulo de geração de grade considerando a disponibilidade docente. A fase de testes e análises será focada na comparação do tempo de alocação manual/legado versus a nova alocação automatizada. A análise de performance do algoritmo Google OR-Tools demonstrará a viabilidade da solução em gerar grades otimizadas em tempo computacional hábil. Os benefícios de uso esperados incluem a drástica redução do trabalho manual da coordenação, a diminuição de conflitos de horário e uma maior satisfação docente. A principal contribuição desta evolução é a flexibilidade e a centralização da informação de disponibilidade diretamente na fonte. **Conclusão:** Espera-se que a evolução desenvolvida seja de alta relevância para a instituição, pois soluciona uma limitação crítica do processo de alocação de disciplinas, especialmente diante de mudanças regulares.

Trilha: Trabalho de Graduação.

GPT Teacher: Desenvolvimento de um Agente de LLM para Programação Assistida em Ambiente VSCode

Rhuan Fernandes, Daniele Oliveira

Universidade Federal de Uberlândia (UFU)
{rhuan.fernandes, danieloliveira}@ufu.br

Introdução: Diante de desafios como desmotivação e dificuldades de compreensão no ensino de programação, a Inteligência Artificial Generativa surge como uma solução promissora. Contudo, a literatura recente alerta que as ferramentas atuais ainda possuem limitações. Focadas na produtividade, essas ferramentas fornecem soluções prontas, fomentando uma “dependência passiva” e inibindo o desenvolvimento do raciocínio lógico e da autonomia do estudante. Esta lacuna define o problema científico: como projetar um agente de IA integrado a um ambiente de desenvolvimento que atue como suporte pedagógico eficaz, em vez de mero assistente de código, a fim de aprimorar ativamente a aprendizagem. **Objetivos:** O objetivo geral é conceber e implementar um protótipo funcional de assistente de IA focado no apoio ao desenvolvimento de competências no aprendizado de programação, priorizando o processo cognitivo do aluno em detrimento da simples produtividade na geração de código. **Método:** Adotou-se a concepção de uma arquitetura de sistema multiagente, implementada como extensão do VsCode. Diferente das soluções focadas em produtividade, a arquitetura proposta desacopla a análise técnica da interação pedagógica por meio de dois agentes especializados: 1) um Agente de Diagnóstico, atuando como desenvolvedor sênior, que analisa o código do aluno e gera um diagnóstico técnico estruturado; e 2) um Agente Tutor, atuando como professor, que recebe esse diagnóstico e o traduz em um diálogo construtivo. Estes agentes são instruídos via engenharia de prompt a guiar o aluno com perguntas para que ele mesmo chegue à solução, evitando respostas diretas e estimulando o pensamento crítico. **Resultados:** A validação foi focada em aspectos funcionais e de desempenho, seus resultados (análise de tempo de inferência e consumo de tokens) evidenciam que a arquitetura proposta é funcionalmente robusta e viável. A ferramenta demonstrou integrar-se ao VSCode, e a abordagem de chat focado no desenvolvimento do raciocínio cumpre o propósito de auxiliar no aprendizado genuíno, validando a viabilidade da hipótese. **Conclusão:** A principal contribuição deste trabalho é a validação de um modelo arquitetural que responde às críticas da literatura sobre assistentes de IA. Enquanto ferramentas convencionais atuam como “oráculos” de informação, a abordagem proposta demonstra a viabilidade de aprimorar as habilidades de ensino de uma LLM via engenharia de prompt, transformando o agente em um facilitador do aprendizado. A integração ao VSCode reduz a carga cognitiva do aluno, alinhando o suporte ao fluxo de trabalho prático. Como limitações, o protótipo analisa apenas arquivos únicos (não projetos completos), a dependência de LLMs de grande porte pode introduzir latência, e o sistema não possui memória de longo prazo para adaptar-se ao histórico do aluno.

Trilha: Trabalho de Graduação.

Implementação e comparação de modelos ML na predição de cepas da Sars-Cov-2

Marcos Paulo G. Pires, Paulo Henrique R. Gabriel

Universidade Federal de Uberlândia (UFU)

{marcos.pires1, phrg}@ufu.br

Introdução: O vírus Sars-Cov-2, pertencente à família dos Coronavírus, é um patógeno altamente infeccioso responsável pela pandemia de Covid-19, ocorrida entre 2020 e 2023. Devido à sua natureza mutável, torna-se essencial o desenvolvimento de estudos sobre possíveis alterações genéticas em seu material genômico, pois essas modificações estão ligadas ao surgimento de novas variantes com diferenças em transmissibilidade, patogenicidade e resposta imunológica. Nesse contexto, a aplicação de modelos de aprendizado de máquina (machine learning), alimentados por bases de dados genéticas, apresenta-se como estratégia promissora para a predição de novas cepas. Essa abordagem permite analisar grandes volumes de dados biológicos e identificar padrões genômicos que indiquem mutações prováveis, contribuindo para o desenvolvimento antecipado de vacinas e para políticas de saúde mais eficazes. Além disso, a comparação entre diferentes algoritmos possibilita avaliar e selecionar aqueles que oferecem maior precisão e desempenho na predição. Assim, este estudo busca investigar, implementar e comparar tais modelos, contribuindo para o avanço de estratégias preditivas aplicadas ao monitoramento genético do Sars-Cov-2 e à mitigação de futuras crises sanitárias. **Objetivos:** O objetivo principal é selecionar e adaptar algoritmos de aprendizado de máquina amplamente utilizados ao contexto de predição de cepas virais do Sars-Cov-2, usando seu código genético como base de análise. Busca-se também comparar e validar o desempenho de cada algoritmo, avaliando sua eficácia e precisão na identificação de padrões genômicos associados a mutações e ao surgimento de novas variantes. **Método:** Serão coletados genomas de diferentes variantes do Sars-Cov-2 por meio de raspagem de dados em bancos públicos, como o GenBank. Em seguida, os dados serão normalizados e filtrados para otimizar o aprendizado das inteligências artificiais. Posteriormente, serão utilizados algoritmos clássicos de Aprendizado de máquina, como Naive Bayes, Random Forest e Decision Tree, adaptados ao contexto de predição viral. Por fim, os modelos serão treinados e avaliados quanto à capacidade preditiva, permitindo um estudo comparativo de desempenho, robustez e precisão. **Resultados:** Espera-se obter um estudo comparativo sobre o desempenho dos algoritmos na predição de novas cepas virais, identificando quais abordagens apresentam melhor capacidade de generalização. **Conclusão:** Embora o estudo ainda esteja em fase de levantamento bibliográfico, os resultados esperados indicam o potencial da pesquisa em ampliar o conhecimento sobre o uso de inteligências artificiais na predição e monitoramento de mutações virais, reforçando sua relevância científica e prática no enfrentamento de futuras ameaças epidemiológicas.

Trilha: Trabalho de Graduação.

Inspeção Heurística de uma Plataforma de Autoria de Jogos Educacionais

Rafael Reis, Rafael Araújo, Renan Cattelan

Universidade Federal de Uberlândia (UFU)

{rafael.reis1, rafael.araujo, renan}@ufu.br

Introdução: Jogos sérios, que são definidos como jogos desenvolvidos com um objetivo principal que não seja unicamente diversão, estão ganhando cada vez mais destaque, seja para treinar, anunciar, simular ou educar a pessoa que estiver jogando em relação a um assunto que o jogo sério se propõe a apresentar. No contexto educacional não é diferente. Jogos têm sido criados como suporte aos processos de ensino e aprendizagem. Previamente, foi construída uma plataforma para gerenciamento de fases de jogos educacionais no contexto da disciplina de Química, uma área que enfrenta grande dificuldade de aprendizado. Esse trabalho realizou uma avaliação da Interação Humano-Computador dessa plataforma.

Objetivos: Esse projeto tem como objetivo o auxílio no desenvolvimento e na validação da ferramenta previamente construída, por meio da condução de uma inspeção heurística para análise de sua usabilidade. **Método:** Foi realizado o procedimento de inspeção de usabilidade, por meio da análise de adequação das telas desenvolvidas frente às 10 heurísticas de Nielsen. A análise foi manual, tela a tela, feita pelo pesquisador e avalia: Visibilidade do estado do sistema; Correspondência entre o sistema e o mundo real; Controle e liberdade do usuário; Consistência e padronização; Reconhecimento em vez de memorização; Flexibilidade e eficiência de uso; Projeto estético e minimalista; Prevenção de erros; Ajuda aos usuários para reconhecerem, diagnosticarem e se recuperarem de erros; Ajuda e documentação. Além das heurísticas, foi avaliado o grau de severidade: Sem importância; Cosmético; Simples; Grave; Catastrófico. Também foi analisada a natureza do problema, se é uma Barreira, um Obstáculo ou um Ruído, a perspectiva do usuário, se o problema é Geral, Preliminar ou Especial e a perspectiva da tarefa, se é um problema Principal ou Secundário. Por fim, foi feita uma descrição do problema, fornecendo contexto, causa, efeito sobre o usuário, efeito sobre a tarefa e correção possível. **Resultados:** Foram encontrados 3 problemas referentes à “Consistência e Padrões” e à “Controle do Usuário e Liberdade”, 2 referentes à “Ajudar os Usuários a Reconhecer, Diagnosticar e Corrigir Erros” e 1 referente à “Flexibilidade e Eficiência de Uso”, “Ajuda e Documentação”, “Prevenção de Erros” e “Estética e Design Minimalista”, totalizando 12 problemas, dos quais 6 eram Cosméticos, 1 era Simples, 2 eram Graves e 3 eram Catastróficos e referente à natureza do problema, 4 eram Barreiras e 8 eram Ruídos. Pela perspectiva do usuário, 9 eram problemas Gerais, 2 eram Preliminares e 1 era Especial e pela perspectiva da tarefa, 4 eram Principais e 8 eram Secundários. **Conclusão:** Mesmo que 50% dos problemas tenham grau de severidade Cosmético, o fato de 25% serem Catastróficos é extremamente preocupante, sendo exigido o reparo desses problemas antes de que a plataforma possa ser disponibilizada para ser usada de fato. Por fim, 17% serem Graves é um ponto de atenção que deve ser analisado com urgência.

Trilha: Trabalho de Graduação.

Integração de Memória Virtual Simulada à Plataforma SimulateOS

Gabriel Reis, Thiago Ribeiro

Universidade Federal de Uberlândia (UFU)

{gabriel.reis1, tpribeiro}@ufu.br

Introdução: O ensino de Sistemas Operacionais (SO) é notório por sua complexidade e abstração, contribuindo para altos índices de reprovação nos cursos de computação, o que motivou o desenvolvimento da Plataforma SimulateOS, um sistema web que simula o comportamento de processos e o gerenciamento de memória de forma visual e interativa. Atualmente, a ferramenta simula a vida de um processo e a memória física, e o objetivo principal é facilitar o entendimento de matérias de SO e arquitetura de computadores, proporcionando uma compreensão essencial de como as operações funcionam “por trás dos panos”. **Objetivos:** O objetivo específico deste projeto é incluir a simulação da memória swap (memória virtual), um mecanismo fundamental que estende a RAM usando o disco rígido e evita travamentos por falta de memória, tornando a visualização interativa do swapping (swap-in e swap-out) um recurso valioso para o aprendizado da gestão dinâmica de memória. **Método:** A metodologia para o aprimoramento seguiu quatro etapas principais: iniciou-se com um Levantamento Bibliográfico sobre simulações e gerenciamento de memória; em seguida, foram definidos os requisitos e realizada a Definição de Requisitos e Modelagem da funcionalidade de swap; a etapa atual consiste no Desenvolvimento do Protótipo, implementando a lógica de movimentação de blocos de memória entre a memória principal e a área de swap para simular o comportamento da memória virtual de forma detalhada; por fim, a fase de Avaliação e Análise de Dados consistirá em analisar os dados coletados nos testes de usabilidade e eficácia da ferramenta. **Resultados:** Como resultados esperados, a Plataforma SimulateOS fornecerá uma ferramenta interativa que permitirá aos alunos visualizar de forma clara a execução dos processos e o gerenciamento da memória, correlacionar a teoria com a prática do sistema simulado, e reduzir a dificuldade de aprendizagem dos conceitos abstratos, impactando positivamente nos índices de aprovação na disciplina, sendo que a implementação da memória swap garante uma simulação mais completa e realisticamente representativa dos mecanismos de um SO moderno. **Conclusão:** Pretende-se que a Plataforma SimulateOS se estabeleça como uma solução de alto potencial para mitigar os desafios de ensino e aprendizagem em Sistemas Operacionais, e sua estrutura visual e interativa tem a capacidade de transformar a absorção do conhecimento técnico, visando a sua consolidação como um recurso educacional de eficácia comprovada após a avaliação do protótipo.

Trilha: Trabalho de Graduação.

JuriZap: Um Assistente Jurídico Brasileiro Gratuito para Leigos baseado em Tecnologias de IA Generativa

Matheus Resende, Fabíola Pereira

Universidade Federal de Uberlândia (UFU)

{matheus.resende1, fabiola.pereira}@ufu.br

Introdução: Em um país onde a democracia é alicerce da convivência social, compreender os próprios direitos é essencial para o exercício pleno da cidadania. No entanto, as leis costumam ser redigidas em uma linguagem técnica e complexa, o que dificulta o entendimento pela maioria da população, bem como restringe a capacidade do cidadão de agir de forma consciente perante a lei. Portanto, urge a criação de uma ferramenta de cunho informacional, a qual torne o conhecimento sobre os direitos do cidadão mais acessíveis e promova a educação jurídica popular. Diante de tal situação, é proposto por este projeto a criação de um assistente jurídico, voltado ao público brasileiro, capaz de interagir com usuários leigos a partir de um modelo de IA. **Objetivos:** Há o objetivo de estudar a modelagem e resolução de problemas com o uso de Modelos de Linguagem de grande escala (LLMs), abordar novas estratégias de “prompts”, aumentar o alcance a informação, disseminar paz e justiça social e diminuir desigualdades, por meio da criação de um aplicativo que permite perguntas sobre problemas jurídicos do usuário. Neste sentido, também é dado como objetivo a diferenciação de outros modelos alternativos a partir da especialização no contexto jurídico brasileiro e simplificação da linguagem usada nas respostas geradas pela Inteligência Artificial, com a finalidade de ser de fácil compreensão para qualquer cidadão. **Método:** A solução será desenvolvida utilizando modelos de linguagem de grande escala (LLMs), em conjunto com a técnica de geração aumentada por recuperação (RAG). Também é previsto a modelagem do aplicativo a partir de uma interface gráfica de interação com o usuário semelhante a um “chat”. Ademais, a literatura existente apresenta avanços significativos em chatbots jurídicos, porém muitos desses modelos carecem de mecanismos em português, bases voltadas para a legislação brasileira, gratuidade e disponibilidade de código aberto para novas contribuições. Em um primeiro instante, é utilizado como base informacional o Código do Consumidor amparado pela legislação brasileira vigente para a geração de respostas mais convincentes sobre processos do âmbito civil e de pequenas causas. **Resultados:** Dentre os resultados iniciais, foi observado que outros modelos conhecidos ainda possuem uma linguagem muito técnica na busca de informações jurídicas, o que pode confundir uma pessoa não especializada no assunto. Outrossim, o modelo proposto apresentou relação mais direta com o código do consumidor brasileiro fornecendo respostas mais precisas e correlatas ao público brasileiro. **Conclusão:** Conclui-se que a formulação de um novo modelo de assistente jurídico voltado ao público leigo é uma alternativa promissora e inovadora, tendo em vista a inexistência de ferramentas voltadas a população geral. Assim, este projeto constitui de uma importante ferramenta social de apoio a educação e informação jurídica no Brasil.

Trilha: Trabalho de Graduação.

Mapeamento Midiático do Viés Racial em Sistemas de Reconhecimento Facial: Resultados Preliminares

Vinícius S. Tadeu, Claudiney R. Tinoco

Universidade Federal de Uberlândia (UFU)
{viniciussilva02t, claudiney.tinoco}@ufu.br

Introdução: O avanço da Inteligência Artificial tem impulsionado a adoção de sistemas de reconhecimento facial na segurança pública. Contudo, essas tecnologias podem carregar vieses algorítmicos que afetam desproporcionalmente a população negra, resultando em prisões injustas e reforçando o racismo estrutural. Diante disso, compreender como a mídia molda o debate público torna-se fundamental para orientar a regulação. Propõe-se, assim, o mapeamento dessa narrativa de forma precisa, filtrando ruídos e superando as limitações de métodos baseados apenas em palavras-chave. **Objetivos:** O principal objetivo é desenvolver e validar um pipeline metodológico de coleta, filtragem semântica e análise de conteúdo jornalístico sobre os impactos do reconhecimento facial na população negra no Brasil (2010–2024). O método busca garantir relevância e consistência do corpus, permitindo análises temporais de sentimentos e de tópicos. Paralelamente, pretende-se superar a baixa precisão das abordagens clássicas, substituindo filtros léxicos por similaridade semântica para obter um conjunto mais qualificado de dados. **Método:** O procedimento metodológico foi dividido em duas fases: (i) na fase de coleta, utilizou-se a API Serper combinando vetores conceituais (Tecnologia, Grupo Social e Impacto), gerando um corpus bruto com 8.582 notícias; e, (ii) na fase de filtragem, para superar a baixa eficiência dos filtros por palavras-chave, foram utilizados embeddings SentenceTransformer, aplicando uma lógica "2 de 3", na qual um título é considerado relevante se for semanticamente semelhante a pelo menos dois vetores-conceito, utilizando um limiar de similaridade de 0,40, estabelecido empiricamente para otimizar a precisão sem perder amplitude. **Resultados:** A filtragem resultou em um corpus preliminar de 5.122 títulos (59,68%). Esse resultado demonstra que a abordagem semântica manteve um volume de dados adequado e aumentou a precisão na seleção de notícias relevantes. A análise temporal sugere que o debate era praticamente inexistente até 2018, mas cresceu fortemente após 2019, atingindo o pico em 2023. A análise de sentimentos (BERT) mostra predominância de cobertura crítica, com 60–65% dos títulos classificados como negativos. **Conclusão:** Os resultados preliminares indicam que o debate midiático é recente, crescente e majoritariamente crítico. A filtragem semântica "2 de 3" com limiar 0,40 mostrou-se uma solução robusta e replicável para a construção de corpus em temas sociotécnicos. Embora a base de dados ainda esteja em fase final de consolidação, o pipeline já demonstra potencial para sustentar análises aprofundadas sobre o viés algorítmico no Brasil e orientar pesquisas futuras.

Trilha: Trabalho de Graduação.

Método Automatizado para Detecção e Mitigação de Ataques Prompt Hacking em LLMs e SLMs

Leandro Rabelo, Diego Molinos

Universidade Federal de Uberlândia (UFU)
{leandro.rabelo, diego.molinos}@ufu.br

Introdução: O avanço dos modelos de linguagem (LLMs e SLMs) tem acelerado a adoção de Inteligência Artificial em múltiplos setores. No Brasil, o percentual de empresas que utilizam IA passou de 16,9% em 2022 para 41,9% em 2024, representando um crescimento de 163,2%. Entretanto, a natureza probabilística desses modelos e a dificuldade de prever comportamentos adversariais mostram-se suscetíveis a ataques cibernéticos, sobretudo o prompt hacking, capaz de contornar restrições internas e gerar vazamento de dados sensíveis. Relatórios recentes de inteligência de ameaças indicam o uso de IA por agentes mal-intencionados para criação de Ransomware-as-a-Service, automação de ataques e fraudes em empregos remotos. A ausência de metodologias automatizadas para detecção e mitigação dessas vulnerabilidades evidencia uma lacuna científica relevante e crescente. **Objetivos:** Esta proposta visa desenvolver uma metodologia automatizada de proteção para LLMs e SLMs contra-ataques de prompt hacking. Para isso, busca-se mapear os principais vetores de ataque que exploram a técnica de prompt hacking, descrevendo seus mecanismos, caracterizando seus padrões e impactos e, por meio de abordagens Red Team e Blue Team, conduzir ataques controlados contra LLMs e SLMs, com ênfase em prompt hacking. **Método:** Trata-se de uma pesquisa exploratória, e o desenvolvimento desta proposta está estruturado em quatro etapas: (a) estudar e compreender os principais vetores de ataque no contexto de prompt hacking para manipular LLMs e SLMs; (b) desenvolver prompts maliciosos com variação semântica e sintática para diferentes tipos de LLMs e SLMs; (c) propor contramedidas para detecção e mitigação desses ataques; e, por fim, (d) validar o método proposto em diferentes modelos de linguagem. **Resultados:** Esta pesquisa encontra-se em fase inicial de desenvolvimento, correspondente à revisão sistemática da literatura com foco em ataques do tipo prompt hacking aplicados a modelos LLM e SLM. Os levantamentos preliminares demonstram escassez de estudos aprofundados que tratem de forma estruturada a detecção e mitigação automatizada dessas ameaças, o que evidencia a relevância e a originalidade da investigação proposta. **Conclusão:** Embora esta pesquisa ainda se encontre em fase inicial e não disponha de resultados consolidados, o tema demonstra elevada relevância científica, tecnológica e social. O crescimento acelerado do uso de LLMs e SLMs em setores críticos amplia a superfície de ataque e expõe vulnerabilidades ainda pouco compreendidas, como o prompt hacking. A ausência de metodologias sistemáticas e automatizadas para detecção e mitigação desses ataques revela um campo de estudo emergente e estratégico, cuja exploração é fundamental para garantir o uso seguro e confiável de LLMs e SLMs. Assim, a presente proposta contribui para o avanço da segurança cibernética aplicada a modelos de linguagem e para o fortalecimento de práticas de defesa cibernética alinhadas aos desafios contemporâneos.

Trilha: Trabalho de Graduação.

Modelo Neuro-Reconfigurável baseado em Modelagem Matemática para Detecção de Ataques DDoS com Baixa Latência

Rafael Vinícius Pimenta, Augusto Silva, Diego Molinos

Universidade Federal de Uberlândia (UFU)

{xrafaelvinicius, augusto.tannus, diego.molinos}@ufu.br

Introdução: Ataques DDoS figuram entre as principais ameaças à segurança da informação. Somente no primeiro semestre de 2025, houve um aumento de 358% deste tipo de ataque, ou seja, são mais de 20,5 milhões de ataques DDoS. Um dos maiores desafios na mitigação deste tipo de ataque reside no tempo de resposta entre a detecção e a ação corretiva. Existem registros de ataques DDoS que levaram apenas 80 segundos para orquestrar 13000 dispositivos IoT, alcançando 5,6 Tbps em um ataque distribuído. Embora as soluções atuais baseadas em Machine Learning (ML) apresentem vantagens na detecção de ataques zero-day frente aos métodos tradicionais, essas abordagens tendem a negligenciar parâmetros de desempenho críticos, como latência de inferência e tempo total de resposta, o que compromete a aplicabilidade dessas soluções em ambientes reais. **Objetivos:** Esta proposta visa o desenvolvimento de uma arquitetura neural reconfigurável em FPGA para detecção em tempo real de ataques DDoS, assegurando baixa latência e alto desempenho em ambientes de alto tráfego. Para viabilizar essa proposta, torna-se essencial racionalizar o modelo de rede neural artificial (ANN) treinado, de modo a representar matematicamente seus parâmetros e operações. Em vez de utilizar ferramentas de síntese de alto nível (HLS) esta proposta irá derivar uma formulação matemática equivalente ao comportamento do modelo treinado, preservando suas características fundamentais. **Método:** Trata-se de uma pesquisa aplicada, de caráter experimental, voltada ao desenvolvimento e à avaliação de uma arquitetura neural reconfigurável para detecção de ataques DDoS. O método proposto compreende seis etapas: (a) Treinamento e teste de uma ANN; (b) Aplicação de XAI para identificação das variáveis independentes e dependentes; (c) Construção de um modelo matemático equivalente à ANN utilizando o método dos mínimos quadrados; (d) Otimização do modelo matemático; e (e) Implementação da arquitetura neural reconfigurável em FPGA. A validação experimental será realizada por meio da análise comparativa entre o modelo e a arquitetura sintetizada. **Resultados:** A proposta encontra-se em andamento, na etapa (a). Até o momento, foram realizadas análises detalhadas das features do dataset, incluindo o balanceamento, a normalização dos dados e a análise de correlações entre as features. **Conclusão:** Espera-se alcançar uma representação fiel do comportamento da ANN treinada, de modo que os testes de inferência validem sua equivalência funcional na detecção de ataques DDoS. Além disso, espera-se comprovar o potencial do uso de FPGAs na área de segurança da informação, destacando sua capacidade de executar modelos inteligentes com baixo tempo de resposta e baixo consumo energético, bem como estabelecer uma metodologia formal e replicável para a implementação de modelos de Inteligência Artificial em arquiteturas reconfiguráveis, contribuindo para o avanço da integração entre IA e hardware em sistemas de defesa cibernética.

Trilha: Trabalho de Graduação.

Modelos Residuais para Classificação de Displasia Oral e Lesões Pulmonares: Um Estudo Comparativo

Luis Felipe G. S. Paim¹, Marcelo Z. Nascimento¹,
Davi Tiago Soares¹, Domingos Lucas L. Oliveira²

¹Universidade Federal de Uberlândia (UFU)

²Instituto Federal de São Paulo (IFSP)

{luis.paim, marcelo.nascimento, davi.tiago}@ufu.br
domingos.latorre@gmail.com

Introdução: A displasia epitelial oral é precursora frequente do carcinoma de células escamosas, um dos tumores mais comuns da cavidade bucal, com transformação maligna entre 6% e 36%. Apesar de avanços diagnósticos e terapêuticos, trata-se de câncer agressivo e ainda muito incidente. O câncer de pulmão é o segundo tipo mais mortal de câncer globalmente, com aproximadamente 1,8 milhões de mortes anuais em 2023. Nesse contexto, algoritmos de processamento de imagens podem aumentar a precisão diagnóstica e apoiar decisões clínicas. **Objetivos:** Investigar algoritmos de análise de imagens para classificar (i) lesões orais em duas categorias (tecido saudável e lesão severa) e (ii) imagens pulmonares em três classes: adenocarcinoma moderadamente diferenciado (aca_md), normal e carcinoma de células escamosas moderado (scc_m). Compararam-se estratégias de inicialização/treinamento e técnicas de data augmentation em ambos os conjuntos com modelos de redes neurais convolucionais residuais. **Método:** Dois conjuntos foram utilizados: 228 imagens histológicas da língua de camundongos e 307 imagens de patologia pulmonar. Avaliaram-se os modelos ResNet18 e ResNet34, variando uso de data augmentation, taxas de aprendizado, número de épocas e dropout. Testaram-se quatro modos: PT-ALL (ajuste de todos os pesos pré-treinados), PT-LB+FC (ajuste do último bloco e da fully connected), PT-FC (apenas fully connected) e FC (treino do zero). Os experimentos foram executados em notebook i7 de 13ª geração, 16 GB de RAM e GPU RTX 3050 (6 GB). **Resultados:** Displasia oral, melhor acurácia (98,55%) com ResNet34 e ResNet18 em PT-ALL, sem augmentation, por 50 épocas. Pior desempenho (81,16%) com ResNet34 em PT-LB+FC combinando augmentation, learning de 0.0001 e 50 épocas. Patologia pulmonar, maior acurácia (98,55%) com ResNet18, sem augmentation, em PT-ALL e taxa de 0,001, o pior resultado ocorreu ao treinar do zero (FC) com augmentation e mesma taxa, tendo 89,86% de acurácia. Modos PT-FC apresentaram desempenhos intermediários, úteis quando a representação final é suficiente, mas, em geral, o fine-tuning completo (PT-ALL) superou as demais configurações. **Conclusão:** Redes neurais convolucionais residual ResNet18 e ResNet34 com transferência de aprendizado e ajuste fino mostraram-se eficazes para análise de lesões displásicas orais e pulmonares. A limitação amostral exigiu estratégias robustas de fine-tuning, o aprendizado transferido foi decisivo para altos desempenhos. Esses achados sugerem que tais sistemas podem tornar o diagnóstico mais objetivo, rápido e preciso na prática clínica.

Trilha: Trabalho de Graduação.

O Impacto do Algoritmo de Grover na Segurança da Criptografia Simétrica

Eduardo Lordão, Giullia Rodrigues

Universidade Federal de Uberlândia (UFU)
{edu.lordao, giullia.rodrigues}@ufu.br

Introdução: A segurança atual é garantida pelos algoritmos criptografados clássicos, o que torna uma maneira extremamente eficaz para proteger dados. Com o avanço da computação quântica e a superposição de bits é possível solucionar problemas de forma mais eficiente. Dessa forma, o Algoritmo de Grover surge como uma ameaça à criptografia clássica. Desenvolvido, não para o objetivo de invadir sistemas, mas para acelerar problema de busca em bases de dados não estruturados, este algoritmo quântico demonstra a capacidade de encontrar um item desejado de forma mais rápida que qualquer outro método clássico. Diante disso, há uma implicação direta na segurança da criptografia simétrica, pois pode reduzir drasticamente o tempo necessário para ataques de Brute Force à chave, colocando em risco a confidencialidade e a integridade de dados protegidos. **Objetivos:** Diante do potencial disruptivo que a computação quântica possui e da iminente ameaça representada pelo Algoritmo de Grover, é necessário analisar suas implicações. Esta pesquisa tem como objetivo principal analisar o impacto do Algoritmo de Grover na segurança da criptografia simétrica, o que contribui para a compreensão das vulnerabilidades apresentadas pela computação quântica, para a conscientização sobre as necessidades de novas abordagens de segurança na era pós quântica e mostrar, por meio de simulações, a eficiência do Algoritmo de Grover sobre outros algoritmos de busca ou Brute Force. **Método:** Inicialmente foi planejada uma pesquisa de natureza teórica e exploratória, utilizando como base predominantemente em revisão bibliográfica e análise conceitual. Porém, conforme o avanço no estudo, foi adicionado uma nova forma de comparar e provar a superioridade do Algoritmo De Grover. Por meio de simuladores quânticos, utilizando um framework de computação quântica (Qiskit), podemos simular a execução de algoritmos quânticos em um computador clássico, focando no número de operações e na probabilidade de sucesso. **Resultados:** Com esta pesquisa, pretende-se provar que o número de operações de algoritmos simétricos cresce em um ritmo diferente em comparação com o Algoritmo de Grover. Dessa forma, espera-se concluir que a busca quântica pode ser uma ameaça substancial para a criptografia simétrica e deve ser mais explorada. **Conclusão:** Espera-se que, com as simulações e análises da complexidade, podemos afirmar que o Algoritmo de Grover quebra a criptografia simétrica, o que impõe uma ameaça crítica ao reduzir significativamente sua segurança. Isso gera uma consequência direta na segurança efetiva e será necessário novas pesquisas para apontar o caminho a seguir para garantir a segurança no futuro da era quântica, provando que, não será necessário o fim da criptografia simétrica, mas precisará ser adaptada. Com isso, a pesquisa irá reforçar a urgência da agilidade criptográfica e a adaptação de sistemas para a computação quântica, uma prioridade imediata para a segurança de dados.

Trilha: Trabalho de Graduação.

Otimização de Consultas em Banco de Dados Astronômico Massivo

Rafael C. Lima¹, Rodrigo C. Boufleur², Altair R. Gomes Júnior¹

¹Universidade Federal de Uberlândia (UFU)

²Associação Laboratório Interinstitucional de e-Astronomia (LIneA)

{rafaeltargaryen99, altair.gomes}@ufu.br

rodrigo.boufleur@linea.org.br

Introdução: A astronomia contemporânea gera volumes massivos de dados. O Portal do Sistema Solar do LIneA lida com uma tabela de predições de ocultações estelares atualmente com 3,5 bilhões de linhas, causando gargalos de performance em consultas temporais e na gestão de dados históricos. Com a era do Big Data. Projetos como o `_Legacy Survey of Space and Time_` (LSST) do Observatório Vera Rubin (Ivezić et al., 2019) e Gaia (Gaia Collaboration, 2023) vem cada vez mais impondo desafios inéditos no gerenciamento de repositórios de dados volumosos. O armazenamento, o acesso e a análise eficiente de catálogos com bilhões de objetos e petabytes de informações se tornaram requisitos fundamentais para o avanço científico. **Objetivos:** A proposta é investigar e comparar estratégias de otimização no PostgreSQL (indexação B-tree/BRIN e particionamento por data), além do arquivamento de dados históricos em formato colunar Apache Parquet, avaliando o impacto dessas estratégias no Django ORM para melhorar o desempenho e a gerenciabilidade da tabela massiva. Objetivos específicos incluem: analisar comparativamente a performance de índices B-tree e BRIN; avaliar o impacto do particionamento por faixa de data; investigar o arquivamento em Parquet, comparando métodos de acesso (FDWs vs PyArrow); e quantificar a sobrecarga do Django ORM versus SQL direto. **Método:** Metodologia baseada em comparações sistemáticas e benchmarks em ambiente controlado. A otimização de dados ativos envolve a comparação entre índices B-tree e BRIN e o uso de particionamento. A gestão de dados históricos testa o arquivamento em Parquet, comparando métodos de acesso (FDWs como `parquet_fdw` ou `duckdb_fdw` e acesso direto via PyArrow). A análise do Django ORM quantifica o overhead comparando consultas ORM com SQL puro. Métricas de desempenho serão coletadas e analisadas estatisticamente para evitar erros de medida. **Resultados:** Espera-se que a combinação de indexação com particionamento demonstre ganhos significativos de performance para consultas temporais na tabela ativa. O arquivamento em formato Apache Parquet deve resultar em alta compressão dos dados. As comparações entre métodos de acesso (FDWs vs. PyArrow) indicarão o melhor equilíbrio entre flexibilidade e performance. A análise do Django ORM deverá quantificar o overhead. **Conclusão:** A conclusão principal será um conjunto de recomendações sólidas e fundamentadas para a otimização do desempenho e gestão de dados do Portal do Sistema Solar. A combinação de técnicas eficientes para dados ativos com estratégias de arquivamento colunar é vital para a sustentabilidade da plataforma do LIneA em um cenário de Big Data.

Trilha: Trabalho de Graduação.

Petúcia e a IA: Estratégia Prática para Desmistificar a Tecnologia na Educação Básica

João Guilherme A. Viana, Thiago Ribeiro, Matheus G. S. Resende

Universidade Federal de Uberlândia (UFU)

{joao.viana1, tpribeiro, matheus-resende}@ufu.br

Introdução: A tecnologia está presente em praticamente todos os aspectos da vida moderna. Os cursos da área de computação têm ganhado destaque e registram procura sem precedentes. Nesse cenário, torna-se essencial apresentar ao ambiente escolar uma visão atualizada do potencial da tecnologia e das diversas áreas que ela abrange, de modo a despertar o interesse dos alunos da Educação Básica. **Objetivos:** Levar às escolas da Educação Básica uma experiência prática e interativa que desperte o interesse dos estudantes pela tecnologia, apresentando, de forma acessível, como diferentes sistemas podem se conectar e atuar em conjunto. Busca-se, ainda, promover uma compreensão mais ampla sobre o papel e as possibilidades do curso de Sistemas de Informação no desenvolvimento de soluções inovadoras. **Método:** Como forma lúdica de exemplificar o uso integrado da tecnologia, está em desenvolvimento uma proposta que demonstra a interação entre diferentes sistemas, permitindo aos alunos observar, de maneira concreta, como esses elementos se conectam e funcionam em conjunto. Será construído um sistema interativo para que os alunos possam “conversar” com uma planta da espécie *Dracaena trifasciata* (Espada-de-São-Jorge), batizada pelo PETS I como “Petúcia”, no qual poderão acessar informações sobre os “sentimentos” da planta, sua necessidade de água (indicada pela umidade do solo), a qualidade do ar, o nível de luminosidade, a temperatura do solo, “borrifadas” virtuais de água para regá-la e dados sobre a espécie e sua biologia, tudo isso integrado a um modelo de Inteligência Artificial. Além disso, cada componente e algoritmo utilizado no sistema será apresentado e explicado em linguagem acessível ao público em geral. O sistema será dividido em cinco módulos: Módulo 1 – Sensores de umidade do solo, umidade do ar, luminosidade, gás e bomba d’água, interligados a um Arduino UNO (ATmega328P); Módulo 2 – Back-end para coletar informações do Arduino; Módulo 3 – Back-end que orquestra os dados coletados, a conversa do usuário e o agente de IA; Módulo 4 – Agente de IA Gemini 2.0 Flash, configurado com informações sobre a espécie e a planta monitorada; Módulo 5 – Interface web para interação com a planta. **Resultados:** Espera-se que o projeto proporcione aos alunos uma experiência prática e significativa, permitindo-lhes compreender como a integração entre hardware, software e Inteligência Artificial pode gerar soluções funcionais. **Conclusão:** O projeto “Petúcia” pretende aproximar os alunos da Educação Básica do universo tecnológico e da Inteligência Artificial, promovendo aprendizagem ativa, interdisciplinar e contextualizada, além de despertar o interesse pelo curso de Sistemas de Informação e evidenciar sua relevância na criação de soluções inovadoras e sustentáveis.

Trilha: Trabalho de Graduação.

Plataforma DebugandoED Repaginada: Integração de Gamificação e Dashboards no Aprendizado de Estruturas de Dados

Guilherme Ventura, Thiago Ribeiro

Universidade Federal de Uberlândia (UFU)

{guilhermehenriqueventura1, tpribeiro}@ufu.br

Introdução: O ensino de Estruturas de Dados, tema fundamental na Computação, frequentemente enfrenta desafios relacionados ao engajamento e à compreensão dos alunos. Nesse contexto, a plataforma DebugandoED foi desenvolvida como uma ferramenta gamificada para auxiliar no aprendizado de estruturas como vetores, pilhas e listas encadeadas. No entanto, a versão original, implementada em PHP, apresenta limitações de usabilidade, desempenho e escalabilidade, dificultando a expansão e a adoção em larga escala. **Objetivos:** Este trabalho propõe a modernização da plataforma por meio da migração para uma arquitetura baseada em Angular no front-end e Node.js no back-end, integrando também um sistema analítico de acompanhamento do desempenho dos alunos. O objetivo geral é melhorar a experiência do usuário, aumentar o engajamento e permitir uma gestão educacional mais eficaz por meio de dados. Como objetivos específicos, destacam-se: a reestruturação da interface para torná-la mais intuitiva e responsiva; a implementação de um sistema de ranking dinâmico e desafios personalizáveis; e a integração de um dashboard analítico para monitoramento em tempo real do progresso discente. **Método:** A metodologia adotada inclui a revisão de trabalhos relacionados, como aqueles que tratam de Learning Analytics Dashboards e gamificação educacional, para embasar as decisões de design e funcionalidade. A solução proposta diferencia-se por combinar gamificação e análise de dados em uma única plataforma, oferecendo tanto motivação lúdica quanto suporte analítico para docentes. **Resultados:** Como resultados parciais, espera-se que a nova versão da DebugandoED apresente maior tempo de sessão e retenção de usuários, com base em indicadores similares aos reportados em estudos anteriores, que apontam aumentos de até 27% no engajamento quando há feedback imediato e elementos gamificados. A avaliação será realizada por meio de testes de usabilidade e análise de métricas de interação, comparando-as com os dados da versão anterior. **Conclusão:** Conclui-se que a modernização da plataforma, aliada à incorporação de mecanismos analíticos, têm potencial para transformar a DebugandoED em uma ferramenta mais robusta e adaptável, contribuindo para o ensino de estruturas de dados de forma mais envolvente, além de permitir a utilização dos dados, para tomar decisões e personalizações de futuras adequações da plataforma.

Trilha: Trabalho de Graduação.

Plataforma Integrada para Detecção e Análise de Ameaças Cibernéticas

Anna Clara Peres, Pedro Miguel Silva, Paulo Henrique Gabriel

Universidade Federal de Uberlândia (UFU)

{annaclararodrigues, miguelsilva@ufu.br, phrg}@ufu.br

Introdução: O crescimento da conectividade digital e da transformação tecnológica nas últimas décadas trouxe consigo o aparecimento de diversas ameaças cibernéticas. A utilização em massa da internet e o avanço de tecnologias, como a Inteligência Artificial (IA), ampliaram ainda mais a variedade de ataque de sistemas corporativos, governamentais e pessoais. Nesse cenário, o Brasil está entre os países mais afetados por incidentes de segurança, com bilhões de tentativas de ataques anuais. Esse contexto reforça a importância de estratégias proativas de defesa e da adoção da Inteligência de Ameaças Cibernéticas (“Cyber Threat Intelligence”, CTI) como um importante instrumento para detecção, mitigação e resposta a incidentes. **Objetivos:** Este trabalho tem como principal objetivo desenvolver uma plataforma integrada para monitoramento e análise de ameaças cibernéticas, capaz de auxiliar profissionais da área de segurança na tomada de decisões proativas. Essa solução busca unir coleta, processamento e visualização de dados em um ecossistema unificado, garantindo praticidade e eficiência no uso da Inteligência de Ameaças Cibernéticas. **Método:** O desenvolvimento da plataforma está sendo conduzido em duas frentes complementares: uma API RESTful e um Dashboard analítico. A API, que será implementada em Java com o framework Spring Boot, será responsável por fornecer dados estruturados provenientes do sistema de detecção e mineração de conteúdo malicioso, especialmente da Dark Web. Já o Dashboard, que foi desenvolvido em Angular, consumirá essas informações e as apresentará por meio de visualizações interativas e filtros dinâmicos. Essa arquitetura promove a integração entre coleta, processamento e interpretação estratégica dos dados. Até o momento, foram levantados os principais requisitos do sistema, desenvolvido o front-end preparado para consumir as APIs futuras e implementado o sistema de autenticação, incluindo autenticação de dois fatores, reforçando a segurança da plataforma. **Resultados:** Como resultado inicial, obteve-se um ambiente funcional para autenticação segura e uma interface intuitiva capaz de exibir informações analíticas assim que a API estiver completamente integrada. A plataforma apresenta potencial para monitoramento de ameaças, configuração de alertas personalizados e identificação de padrões e tendências em ameaças cibernéticas. **Conclusão:** Esta plataforma deve contribuir diretamente para o fortalecimento da segurança cibernética em ambientes institucionais e corporativos. Com a conclusão da implementação da API, espera-se alcançar uma solução completa de suporte à análise de ameaças, fornecendo dados estratégicos que auxiliem profissionais de segurança na tomada de decisões proativas e embasadas.

Trilha: Trabalho de Graduação.

Proposta de Implantação de um Chatbot para Avaliação no impacto das Operações do HC-UFU/EBSERH

Jonathan D. Aquino, Daniel Stefany D. Caetano

Universidade Federal de Uberlândia (UFU)
{jonathandavi, daniel.caetano}@ufu.br

Introdução: A crescente adoção de soluções de automação conversacional tem impulsionado o desenvolvimento de bots voltados a agilizar interações, padronizar atendimentos e integrar serviços digitais. Inserido nesse contexto, surgiu a proposta de um chatbot para o Hospital de Clínicas da UFU (HC-UFU), com o objetivo de modernizar a comunicação e reduzir dependência de canais tradicionais (sites e atendimentos manuais), que frequentemente tornam o fluxo de informações mais lento e menos confiável. **Objetivos:** Os objetivos principais do projeto giram em torno de ampliar a capacidade de atendimento aos usuários, automatizar tarefas repetitivas, oferecer suporte instantâneo, responder a uma ampla variedade de perguntas e realizar ações simples, como fornecer informações, links, agendamentos, direcionamentos a sítios e portais, entre outros, possibilitar a disponibilidade de atendimento automatizado 24 horas por dia, ampliar a experiência do usuário, reduzir custos operacionais e implementar novas tecnologias no HC-UFU. **Método:** Para a versão inicial do sistema, são considerados os seguintes requisitos funcionais e técnicos: 1-Definição de áreas de interesses como, hospital e capelania; 2-Definição de fluxos separados para ambos, com comandos e gatilhos para a validação de mensagens; 3-Estrutura lógica condicional para comunicação entre o usuário e o sistema; 4-Implementação utilizando a biblioteca whatsapp-web-js feita em Node.js. 5-Scan de QRCode para que os usuários possam iniciar a conversa com o sistema sem dificuldades. **Resultados:** Até o momento, foram estruturados os componentes centrais do bot, incluindo o núcleo de roteamento de mensagens e orquestração de comandos para o recebimento e resposta a mensagens dos usuários, camadas utilitárias para formatação e validação onde necessário. Grande parte do segmento da capelania pronto, o usuário pode desde solicitar uma visita de um capelão com o bot até pedir informações simples para suas dúvidas, todos os dados necessários estão sendo armazenados em banco de dados utilizando MySQL. No segmento do hospital o bot já é capaz de responder perguntas como, qual o horário de visita? Como eu chego no hospital? entre muitas outras e redirecionar os usuários para os sites oficiais do HC-UFU. **Conclusão:** As próximas etapas incluem finalizar ambos os fluxos seguindo a lógica de resposta, implementar uma central de controle para que os funcionários e técnicos do HC-UFU possam ter controle total do bot, podendo alterar mensagens, consultar dados etc. Em um futuro a ideia é adicionar integração com inteligência artificial (IA), para que o bot se desprenda das mensagens prontas e passe a ter mais funcionalidades podendo fazer mais procedimentos de atendimento dentro do hospital.

Trilha: Trabalho de Graduação.

Proposta de Intervenções Automatizadas via Chatbot na Plataforma DebugandoED

João Guilherme A. Viana, Thiago P. Ribeiro

Universidade Federal de Uberlândia (UFU)

{joao.viana1, tpribeiro}@ufu.br

Introdução: Atualmente, o uso da tecnologia voltada para o aprendizado tem se intensificado significativamente, com destaque para ferramentas especializadas em instrução de conteúdos específicos e, especialmente, aquelas que utilizam recursos de Inteligência Artificial. Segundo estudos, estas ferramentas auxiliam consideravelmente o aprendizado de conteúdos com maior interatividade e personalização. No entanto, ainda existem obstáculos de adaptação e compreensão por parte dos alunos, motivo pelo qual a IA pode atuar como agente facilitador e apoio no processo de ensino-aprendizagem. **Objetivos:** O principal objetivo dessa pesquisa é dispersar estes obstáculos, utilizando a Inteligência Artificial como direcionador para que o aluno compreenda seus pontos fracos sobre determinado assunto abordado que, neste caso, refere-se à plataforma DebugandoED, destinada a auxiliar no ensino de Estrutura de Dados. Assim, o recurso abordado nesta plataforma terá como papel identificar o que levou o aluno a responder incorretamente o que foi solicitado na questão, bem como direcioná-lo sobre como chegar ao resultado correto. **Método:** Esta ferramenta aborda o conceito de implantação de Redes Neurais Artificiais (RNA) e Aprendizado de Máquina (Machine Learning). Ao aluno responder uma sentença na plataforma, será feita uma validação do resultado: caso tenha acertado, nada será feito e prosseguirá para as próximas atividades; caso contrário, a rede entra em ação para auxiliar o usuário a compreender onde errou e o que pode ter levado a isso. Para que a RNA saiba como direcionar o aluno diante de um erro, são necessários dados importantes para que a rede consiga ligar os pontos e entregar uma resposta mais precisa; assim, ela precisa de uma base de dados com erros anteriores do usuário, a fim de correlacionar os pontos fracos e de baixa adesão que o aluno apresenta diante de um conteúdo X ou Y e verificar se isso interfere em resultados gerais. Também é preciso um conjunto de dados vindo da sentença apresentada, indicando a resposta correta, a alternativa mais provável de ser escolhida e possíveis fatores de confusão. A partir dessa análise, a RN, integrada a um chatbot na plataforma, poderá compreender explicitamente a razão pela qual o aluno não está conseguindo chegar ao resultado esperado. **Resultados:** Espera-se que a RNA consiga interpretar e gerar respostas condizentes e com acurácia favorável a partir dos dados obtidos. **Conclusão:** Espera-se que a integração da Inteligência Artificial à plataforma DebugandoED contribua significativamente para superar as dificuldades de aprendizado em Estrutura de Dados. Desse modo, prevê-se que o sistema identifique falhas de compreensão, ofereça direcionamentos personalizados e promova um aprendizado mais eficaz, autônomo e adaptado às necessidades de cada aluno.

Trilha: Trabalho de Graduação.

Recomendação de Formação de Equipes de Avaliação para o PNLD com Foco em Equidade e Representatividade

Pablo Araújo, Rafael Araújo

Universidade Federal de Uberlândia (UFU)

{rcapablo7, rafael.araujo}@ufu.br

Introdução: O Programa Nacional do Livro e do Material Didático (PNLD) é uma política pública executada pelo FNDE e pelo MEC, destinada a disponibilizar obras didáticas, pedagógicas e literárias de forma sistemática e gratuita. O programa compreende um ciclo com várias etapas, dentre as quais a avaliação pedagógica se destaca por garantir que os materiais estejam alinhados à legislação e às concepções pedagógicas do país. Nessa etapa, profissionais da educação de todo o Brasil, pertencentes a um banco de avaliadores, analisam as obras com base em um material guia. Dentro disso, a seleção dos avaliadores para a criação de grupos é essencial para assegurar uma avaliação plural e coesa. No entanto, essa etapa pode se tornar um gargalo devido ao grande número de avaliadores e à complexidade dos critérios de formação. O banco de avaliadores do PNLD já soma cerca de 14.500 docentes, embora esse número seja muito positivo, ele intensifica o desafio de compor equipes qualificadas e heterogêneas, que representem de forma equilibrada a diversidade regional, de gênero, raça e outras dimensões sociais relevantes. **Objetivos:** Neste cenário, o objetivo desta pesquisa é criar um método de formação de grupos mais equitativos e heterogêneos, analisando tecnologias utilizadas em processos de formação de equipes. Considera-se heterogêneo o grupo que apresenta uma distribuição equilibrada de diferentes perfis, contemplando critérios como formação acadêmica, diversidade demográfica, distribuição geográfica, inclusão e trajetória educacional. O projeto justifica-se pela necessidade de apoiar políticas públicas com ferramentas que ampliem a transparência e a eficiência da seleção de avaliadores, contribuindo para um ambiente educacional mais justo e representativo. **Método:** A pesquisa segue uma abordagem exploratória, buscando compreender o processo de avaliação do PNLD para formalizar restrições e objetivos de otimização, além de investigar tecnologias adequadas ao desenvolvimento da solução computacional. A prova de conceito será desenvolvida com um conjunto de dados sintéticos que reproduz a estrutura dos dados reais dos avaliadores, permitindo simular cenários diversos de formação de equipes. **Resultados:** A validação da prova de conceito será feita por meio dos dados simulados. O desempenho do modelo será comparado ao método atual de seleção, utilizando métricas e índices de diversidade. A hipótese é que o modelo proposto formará equipes com diversidade significativamente superior e com maior eficiência que o processo tradicional. A interface de visualização permitirá uma análise interativa dos resultados, reforçando a sua viabilidade. **Conclusão:** O trabalho propõe uma solução computacional para um desafio de gestão em políticas públicas. A principal vantagem da abordagem está na sua capacidade de equilibrar critérios de forma flexível e rápida.

Trilha: Trabalho de Graduação.

Reconhecimento de Emoções em Sinais Sonoros: Uma Revisão de Modelos e Metodologias

Gabriel Vieira, Marcelo Z. Nascimento

Universidade Federal de Uberlândia (UFU)
{g.antonio, marcelo.nascimento}@ufu.br

Introdução: O campo de reconhecimento de emoções no meio computacional oferece grande destaque quando é avaliado o desenvolvimento de técnicas que se propõem a otimizar a comunicação e interação entre pessoas e sistemas. O reconhecimento de emoções por meio de sinais de áudio (do inglês, Speech Emotion Recognition (SER)) apresenta um tema relevante, motivando o desenvolvimento de estudos voltados à criação de modelos em busca da garantia de melhor desempenho em tarefas de reconhecimento. Entre as abordagens em SER, predominam modelos que categorizam emoções em um conjunto finito de rótulos, favorecendo a leitura e a interpretação dos resultados e sua aplicação em cenários como sistemas de feedback, centrais de atendimento, segurança e pânico, perícia e outros. Diante das diversas aplicações e técnicas possíveis, surge a necessidade de realizar um estudo que se desdobre sobre a análise do estado da arte das produções sobre o tema, a fim de compreender quais são as principais características de procedimentos de reconhecimento de emoções em áudio voltados à classificação, e as principais contribuições do tema para os contextos onde se aplica. **Objetivos:** O presente trabalho busca estabelecer uma compreensão sobre o estado da arte de produções em SER, bem como fornecer uma análise sobre as principais técnicas utilizadas através da implementação de modelos descritos na literatura e a avaliação de sua acurácia no que diz respeito a tarefas de classificação, apresentando uma visão geral dos principais desafios e limitações das técnicas documentadas. **Método:** Realizou-se uma revisão sistemática da literatura que embasou a implementação de técnicas e modelos. A linguagem Python foi utilizada para a construção dos modelos e o software openSMILE para extração de características de áudio. **Resultados:** A revisão sistemática selecionou 10 estudos que permitiram experimentos com métodos de classificação tradicionais em aprendizado de máquina sobre o conjunto de dados RAVDESS, produzindo métricas próximas às relatadas no estado da arte. **Conclusão:** Foram identificadas as etapas centrais de um fluxo de SER e as principais dificuldades para aprimorar o desempenho. Os pontos críticos concentram-se na seleção de características relevantes do sinal e na representação do áudio para análise computacional. As técnicas atuais enfrentam entraves sobretudo na definição de descritores capazes de capturar informações suficientes para distinguir padrões emocionais complexos na fala.

Trilha: Trabalho de Graduação.

Redesign Centrado no Usuário: Uma Intervenção no Site das Bibliotecas da UFU sob a Perspectiva de Interação Humano Computador

Nicolly Luz, Rogério Filho, Daniel Caetano, Otávio Côrtes

Universidade Federal de Uberlândia (UFU)

{nicolly.luz, rogerio.filho, daniel.caetano, otavio.cortes}@ufu.br

Introdução: O presente trabalho foi concebido com o objetivo de aplicar os princípios da Interação Humano-Computador (IHC) na análise crítica e no redesign de uma interface real. O estudo de caso concentrou-se no site das bibliotecas da Universidade Federal de Uberlândia (UFU), uma plataforma de alta relevância para a comunidade acadêmica, porém identificada com significativas falhas de usabilidade e acessibilidade. Desta forma, o projeto buscou compreender as limitações existentes, propor soluções centradas no usuário e, por meio de métodos de avaliação consolidados, verificar a contribuição das mudanças implementadas para uma experiência mais intuitiva e eficiente. **Objetivos:** O objetivo central foi aprimorar a usabilidade e a acessibilidade do site, buscando uma experiência de usuário mais funcional e satisfatória. Especificamente, o projeto focou em: identificar problemas críticos de design e navegação; propor melhorias baseadas nas boas práticas de IHC; desenvolver protótipos de alta fidelidade; e avaliar a eficácia das soluções por meio das Heurísticas de Nielsen. **Método:** A metodologia adotada foi estruturada em um ciclo de quatro etapas interligadas: análise, síntese, intervenção e avaliação. Inicialmente, a fase de análise combinou a aplicação de questionários para mapear a percepção dos usuários com uma investigação de natureza quali-quantitativa aprofundada. Para a síntese, foram empregadas técnicas consolidadas, como a Análise Hierárquica de Tarefas (HTA), o Mapa de Jornada do Usuário e o Card Sorting, permitindo uma compreensão detalhada dos fluxos de interação, dos pontos de atrito e do modelo mental dos usuários. Na etapa de intervenção, as telas originais foram integralmente reformuladas, priorizando a clareza, a consistência visual e os requisitos de acessibilidade. Por fim, a avaliação foi conduzida mediante a aplicação das dez Heurísticas de Usabilidade de Nielsen aos protótipos desenvolvidos, validando teoricamente a eficácia das melhorias propostas. **Resultados:** A análise inicial revelou problemas como navegação complexa, excesso de informações e baixa integração funcional. Em resposta, foram desenvolvidas novas versões para telas-chave (Home, Login, Busca, Serviços, etc.), concentrando-se na simplificação dos menus laterais, adoção de linguagem clara, design minimalista e coerência visual. Além disso, a inclusão de uma página dedicada à acessibilidade e de uma central de ajuda interativa foi fundamental para promover maior autonomia e inclusão ao usuário. **Conclusão:** O projeto validou o impacto positivo do redesign na interface das bibliotecas da UFU, tornando-a significativamente mais acessível, intuitiva e aderente às boas práticas da Interação Humano-Computador (IHC), assim como das Heurísticas de Nielsen. A aplicação rigorosa das técnicas de análise e prototipagem foi crucial para alinhar o sistema às necessidades reais dos usuários, resultando em uma experiência de navegação notavelmente mais fluida e eficiente.

Trilha: Trabalho de Graduação.

Robustez ao Aumento de Dados: Impacto do Aumento de Dados nas Redes VGG e SqueezeNet na Classificação de Lesões Histológicas

Luana Borges, Vitória Silva, Domingos Latorre, Marcelo Nascimento

Universidade Federal de Uberlândia (UFU)

{luana.borges1, vitoria.fernandes}@ufu.br

{domingos.latorre, marcelo.nascimento}@ufu.br

Introdução: A análise de imagens histológicas é fundamental para o diagnóstico de diversas doenças. Em muitos contextos, os conjuntos de dados são limitados ou apresentam distribuição desigual entre classes, o que prejudica o desempenho de modelos de redes neurais convolucionais (CNNs), ferramentas que podem auxiliar patologistas em seus diagnósticos. O estudo de estratégia de aumento de dados surge como uma alternativa para ampliar a variabilidade das imagens e tornar os modelos mais robustos. No entanto, o impacto dessas técnicas depende da arquitetura CNN utilizada e da capacidade de preservação das características morfológicas relevantes das amostras. **Objetivos:** Este trabalho investiga como diferentes estratégias de aumento de dados influenciam o desempenho dos modelos VGG13, VGG19 e SqueezeNet na classificação de lesões histológicas. As arquiteturas VGG13 e VGG19 foram selecionadas por apresentarem múltiplas camadas convolucionais com filtros de tamanho reduzido, o que favorece a extração hierárquica de características como texturas, bordas e estruturas celulares. Já o modelo SqueezeNet foi escolhido por sua arquitetura compacta, capaz de reduzir significativamente o número de parâmetros sem comprometer o desempenho. **Método:** As CNNs foram treinadas utilizando quatro estratégias: treinamento do zero (fFrom Scratch), Fine-Tuning do classificador (PT-FC), Fine-Tuning com descongelamento parcial (PT+LB-FC) e Fine-Tuning completo (PT-ALL). Foram testadas diversas combinações de transformações, incluindo recortes aleatórios, inversões, ajustes de cor, borramento e rotações. Para comparação detalhada, analisou-se o efeito do aumento de dados nas CNNs e a preservação das estruturas celulares e das arquiteturas teciduais. Os experimentos foram conduzidos em dois cenários: um dataset de displasia oral balanceado entre classes saudável e severa, e um dataset pulmonar com três classes desbalanceadas. **Resultados:** No dataset de displasia oral, a melhor configuração incluiu RandomResizedCrop, inversão horizontal e RandomErasing, o que aumentou a variabilidade sem alterar padrões diagnósticos, resultando em maior acurácia e menor overfitting. Transformações mais intensas, como rotações e ajustes de cor (ColorJitter), reduziram o desempenho ao modificar a orientação e a coloração dos tecidos, interferindo na detecção de características morfológicas importantes. No dataset pulmonar, o uso moderado de ColorJitter e a aplicação ocasional de desfoque melhoraram a generalização em um cenário desbalanceado, enquanto as rotações novamente deterioraram o desempenho. **Conclusão:** Os resultados demonstraram que as técnicas de aumento de dados melhoram o desempenho dos modelos quando preservam a organização histológica essencial. Transformações moderadas e compatíveis com a variabilidade real do tecido geram modelos mais robustos, enquanto alterações que distorcem as estruturas celulares tendem a prejudicar a classificação.

Trilha: Trabalho de Graduação.

Seleção de Vértices Semiautomática e Reconstrução Automática de Malhas de Cabeças 3D para Fins de Animação

Víctor Senema, Daniel Abdala

Universidade Federal de Uberlândia (UFU)

{victor.senema, abdala}@ufu.br

Introdução: A modelagem tridimensional sempre foi uma das áreas de interesse na computação. Inicialmente aplicada na engenharia e, posteriormente, expandida para o entretenimento. Na área de modelagem, os modelos 3D são formados por arestas, vértices e faces, que juntos formam a malha do objeto. Existem diferentes tipos de modelos, alguns usados para a renderização de alta qualidade e outros para animações. Os modelos utilizados para renderização priorizam o aspecto visual e o volume aparente. É como se fossem esculturas de argila, nas quais se adiciona material para dar o volume desejado e, posteriormente, se retira o excesso e se adicionam os detalhes. Nesse caso, a organização da malha não é muito relevante. Por outro lado, os modelos destinados à animação exigem uma construção mais cuidadosa, além de apresentarem as mesmas qualidades dos modelos exclusivos para renderização também devem ser movidos sem apresentar deformações indesejadas, necessitando de uma topologia bem estruturada. O objetivo do trabalho é criar um plugin para o Blender que tome como entrada uma malha para renderização e alguns pontos de controle selecionados pelo usuário e então produza como saída um modelo apto a ser animado. **Objetivos:** o trabalho tem como foco desenvolver uma ferramenta para o Blender que permita reconstruir a topologia de modelos 3D de cabeça de forma semiautomática. O plugin receberá de entrada uma malha inadequada para animação, além de pontos de controles definidos pelo usuário, e retornará uma malha adequada para realizar deformações no processo de animação. **Método:** O desenvolvimento será realizado utilizando API do Blender em Python. Inicialmente o usuário selecionará pontos críticos para que seja possível criar um esqueleto topológico o qual servirá para identificar largura, direções e conectividades. A clusterização hierárquica será a próxima etapa, na qual a malha será subdividida visando princípios da percepção humana. E por fim será reconstruída utilizando algoritmos como de reconstrução de malha por método de Poisson. Por fim serão utilizados modelos de diferentes densidades para validar a ferramenta. **Resultados:** Testes preliminares demonstraram que a plataforma do Blender oferece suporte adequado às operações necessárias, como estruturas de dados e métodos para manipulação de malha. É esperado que o plugin produza malhas reorganizadas de forma consistente, evitando problemas como descontinuidade de superfície e dobras incorretas. O protocolo de validação contemplará métricas como quantidade de vértices, qualidade das deformações pós processo de animação e fidelidade ao modelo original. **Conclusão:**

Trilha: Trabalho de Graduação.

Sistema de Comparação de Preços para E-Commerce com Machine Learning

Guilherme Machado, Fernanda Santos

Universidade Federal de Uberlândia (UFU)

{guiguicastilho34, fmcsantos}@ufu.br

Introdução: O comércio eletrônico brasileiro experimentou um crescimento exponencial, impulsionado pela pandemia da COVID-19, ampliando o número de lojas virtuais e a diversidade de produtos. Esse cenário trouxe como desafio a grande variação de preços para um mesmo item em plataformas como Amazon, Mercado Livre e Shopee, agravada pela heterogeneidade semântica e a falta de padronização nos nomes dos produtos. Para enfrentar esse problema, este trabalho propõe um sistema de comparação de preços utilizando Web Scraping, Processamento de Linguagem Natural (PLN) e Aprendizado de Máquina (ML) para coletar, padronizar e comparar itens de diferentes e-commerces. **Objetivos:** O objetivo principal é desenvolver um sistema computacional para comparação de preços em tempo real, focado no e-commerce brasileiro. Os objetivos específicos são: 1) Implementar rotinas de Web Crawling e Scraping para 11 grandes e-commerces, visando a extração de atributos-chave (nome, preço, URL); 2) Aplicar técnicas de PLN para a normalização textual dos dados coletados, resolvendo o problema da não-padronização; 3) Validar a eficácia de modelos de ML (KNN, Random Forest) na tarefa de correspondência de produtos (entity matching), garantindo que itens semanticamente equivalentes sejam agrupados; 4) Criar uma interface web intuitiva (Django) para apresentação dos resultados. **Método:** A metodologia segue um fluxo de prototipagem experimental. A coleta de dados é feita com a biblioteca BeautifulSoup. Em seguida, bibliotecas como NLTK e spaCy realizam a normalização textual (tokenização, lematização, remoção de stopwords). Posteriormente, modelos de ML são treinados com Scikit-learn para identificar similaridades entre os produtos e permitir comparações precisas. Os resultados processados são armazenados em banco de dados e exibidos ao usuário por meio da interface Django. **Resultados:** O sistema está em desenvolvimento, com módulos essenciais concluídos e validados. Os componentes de Web Scraping estão operacionais e funcionais para os 11 e-commerces. A arquitetura de back-end (Django) e os modelos de banco de dados estão implementados. A interface de busca e exibição de resultados já permite ao usuário consultar produtos e visualizar os dados brutos coletados. A etapa atual concentra-se no treinamento e ajuste dos modelos Random Forest e KNN para a classificação dos itens. **Conclusão:** O trabalho demonstra a viabilidade de um sistema inteligente capaz de comparar preços de forma eficiente e adaptado ao contexto brasileiro. A principal contribuição é a criação de uma ferramenta que resolve o problema central da correspondência de produtos através de IA. Ao ser concluído, o sistema oferecerá economia de tempo, aumento da precisão nas comparações e maior transparência ao consumidor.

Trilha: Trabalho de Graduação.

Sistema de Recomendação baseado em Conteúdo para Programação Competitiva: Uma Abordagem de Personalização de Treino

Júlia Koba, Daniele Oliveira

Universidade Federal de Uberlândia (UFU)
{julia.koba, danieloliveira}@ufu.br

Introdução: A programação competitiva é um esporte mental que consiste em resolver problemas lógicos e matemáticos, avaliando-se por número de problemas resolvidos, tempo e eficiência. A falta de orientação personalizada e a indisponibilidade de treinadores geram alta desistência entre estudantes. Para contornar isso, foi implementado um sistema de recomendação de exercícios, planos de estudo e avaliação de estatísticas do usuário. Observa-se escassez de pesquisas que integrem inteligência artificial e recomendação educacional especificamente nesse contexto, limitando soluções personalizadas e escaláveis. **Objetivos:** Aumentar a motivação e aumentar a eficiência do aprendizado em programação competitiva por meio de um sistema de recomendação inteligente, capaz de adaptar-se continuamente ao progresso e ao perfil do estudante. Busca-se desenvolver um ambiente personalizado que identifique lacunas de conhecimento, proponha exercícios adequados ao seu nível e auxilie na estruturação de trilhas de estudos evolutivas. Além disso, tem como objetivo avaliar o impacto de técnicas de IA e modelos de linguagem (LLMs) na geração de recomendações mais precisas e motivadoras, alinhadas ao desempenho e aos objetivos individuais dos usuários. **Método:** Adotou-se abordagem baseada em conteúdo. O sistema analisa categorias e dificuldade dos problemas e as compara com o perfil do usuário, combinando dados dinâmicos (histórico de submissões do Codeforces) e dados estáticos declarados (período da faculdade, estilo de aprendizado e tópicos de interesse). A LLM realiza 1) análise semântica das descrições para inferir categorias faltantes e 2) geração de roteiros de estudo adaptados ao perfil. As ferramentas utilizadas foram Spring Boot no backend, React.js no frontend e PostgreSQL, integrando a API do Codeforces. **Resultados:** O protótipo foi validado com “handles” (apelidos) do Codeforces iniciantes, intermediários e avançados. As recomendações refletiram coerência com as categorias inferidas, identificando tópicos recorrentes e sugerindo exercícios progressivos. Em comparação com abordagens anteriores baseadas em filtragem colaborativa ou análise sequencial, a estratégia aqui combina conteúdo e análise semântica, com foco no perfil do aluno e planos pedagógicos adaptativos, ampliando relevância e diversidade das recomendações. **Conclusão:** O estudo confirma a viabilidade de um sistema de recomendação baseado em conteúdo e IA para programação competitiva. Destaca-se a personalização dupla: adaptação ao desempenho real e aos objetivos declarados. O sistema atua como “tutor virtual”, mitigando desmotivação e promovendo autonomia. Limitações incluem dependência de APIs externas e risco de “bolha de filtro”. Futuramente, pretende-se evoluir para abordagem híbrida com filtragem colaborativa para ampliar descoberta de novos tópicos e diversidade de aprendizado.

Trilha: Trabalho de Graduação.

Tecnologia Acessível: Uso de VANTs de Baixo Custo no Monitoramento Ambiental e Ocorrências Críticas

Victor Brizante, Diego Molinos

Universidade Federal de Uberlândia (UFU)
{victor.brizante, diego.molinos}@ufu.br

Introdução: O monitoramento ambiental e a resposta rápida a desastres naturais ou acidentes industriais são desafios logísticos e financeiros. Veículos Aéreos Não Tripulados (VANTs) destacam-se como ferramentas eficazes na coleta de dados em tempo real, reduzindo riscos e custos operacionais. No monitoramento ambiental, permitem identificar focos de calor, mapear vegetação e medir áreas afetadas com alto detalhamento. Essas informações subsidiam ações de combate a incêndios, avaliação de impactos e planejamento de recuperação. Entretanto, o alto custo de sistemas comerciais limita o acesso a essa tecnologia, especialmente em projetos com recursos reduzidos. O desenvolvimento de VANTs acessíveis democratiza o uso dessa ferramenta, ampliando sua aplicação em pesquisas, gestão pública e iniciativas comunitárias, além de promover inovação aberta e sustentabilidade.

Objetivos: Desenvolver e dominar todo o ciclo de manufatura de VANTs de baixo custo, desde o projeto até a validação em campo. O trabalho propõe uma arquitetura de hardware e software modular, de fácil replicação e manutenção. Também busca implementar uma interface de controle e telemetria compatível com futuras expansões, como GPS e comunicação sem fio de longo alcance, validando a solução como ferramenta robusta e acessível para monitoramento ambiental e resposta a ocorrências críticas.

Método: O método abrange o projeto e a produção do VANT, priorizando materiais acessíveis e sustentáveis, como peças impressas em 3D (PLA/PETG), isopor reutilizado e componentes eletrônicos de prateleira (COTS). O sistema de controle utiliza plataformas Arduino e ESP32 com comunicação RF 2.4 GHz, garantindo confiabilidade e compatibilidade. As etapas incluem: (a) construção e testes estruturais; (b) desenvolvimento da controladora de voo open source; (c) integração de sensores e câmeras; e (d) validação em cenários simulados de monitoramento. O software foi desenvolvido de forma modular, permitindo ajustes de estabilidade, calibração de atuadores e adaptação a diferentes configurações aerodinâmicas.

Resultados: Foram construídos múltiplos protótipos funcionais com os materiais descritos. Testes de campo validaram a estabilidade e robustez do enlace de comunicação entre transmissor e receptor, demonstrando controle confiável em linha de visada. Também foram coletados dados de telemetria (tensão, corrente e sinal RF), utilizados para otimizar o desempenho elétrico e ampliar a autonomia de voo.

Conclusão: Os resultados comprovam a viabilidade técnica e econômica da arquitetura proposta. As próximas etapas incluem a integração de sensores e câmeras para captura de dados ambientais e o desenvolvimento do voo autônomo. O sistema resultante visa constituir uma plataforma aberta e replicável, capaz de apoiar projetos de monitoramento, pesquisa e gestão de riscos em diferentes contextos.

Trilha: Trabalho de Graduação.

Um Estudo Comparativo entre Sistemas de Gerenciamento de Bancos de Dados de Vetores com Foco em Dados Textuais

Nicolly Luz, Fabíola Pereira

Universidade Federal de Uberlândia (UFU)
{nicolly.luz, fabiola.pereira}@ufu.br

Introdução: O avanço acelerado da inteligência artificial e, em especial, dos modelos de linguagem de grande escala (LLMs), tem impulsionado a necessidade de soluções eficientes para o armazenamento e recuperação de informações complexas. Nesse contexto, os bancos de dados de vetores surgem como uma tecnologia essencial, permitindo o armazenamento, indexação e busca de representações vetoriais comumente conhecidas como embeddings, amplamente utilizadas em aplicações de IA e em sistemas de Retrieval-Augmented Generation (RAG). Apesar da crescente adoção desses sistemas, ainda há lacunas na literatura quanto à análise comparativa de desempenho, escalabilidade e eficiência entre as diferentes ferramentas disponíveis. **Objetivos:** Este trabalho tem como objetivo principal realizar um estudo comparativo entre Sistemas Gerenciadores de Bancos de Dados de Vetores (SGBDVs), avaliando aspectos qualitativos e quantitativos, com enfoque na medição de eficiência na recuperação de informações vetoriais e no tempo de retorno dos resultados, respectivamente. Motivado pela variedade de ferramentas disponíveis (como pgvector, Cassandra, MongoDB, Chroma DB, Qdrant, Weaviate, entre outros.) e pela ausência de análises sistemáticas sobre suas vantagens e limitações, o estudo busca compreender o comportamento dessas ferramentas em cenários que exigem alto desempenho em consultas vetoriais, os objetivos específicos incluem: (i) selecionar os SGBDVs a serem comparados; (ii) construir uma base vetorial a partir de textos; (iii) propor métricas que assegurem uma análise justa e representativa; e (iv) implementar experimentos e avaliar os resultados obtidos. **Método:** A metodologia adotada compreende as seguintes etapas: estudo teórico sobre o funcionamento e as arquiteturas de bancos de dados de vetores; seleção das ferramentas a serem comparadas; definição e preparação dos dados textuais que gerarão os vetores; escolha das métricas de avaliação; implementação de casos de teste; e, por fim, coleta e análise dos resultados experimentais. **Resultados:** Embora os resultados ainda estejam em fase de coleta, já foi possível observar diferenças significativas nos tempos de consulta entre os diferentes SGBDVs avaliados. Tais variações sugerem que cada ferramenta apresenta particularidades estruturais que impactam diretamente sua eficiência e aplicabilidade em determinados contextos. **Conclusão:** Espera-se, ao final da pesquisa, fornecer uma análise comparativa abrangente que contribua para a escolha informada de tecnologias em projetos que envolvam grandes volumes de dados vetoriais. Além disso, os resultados poderão servir de base para futuras otimizações e aprimoramentos na integração entre bancos de dados e sistemas de IA, evidenciando a relevância deste estudo para o avanço da área.

Trilha: Trabalho de Graduação.

Um Sistema Tutor Inteligente para Apoio ao Ensino de Computação

Alexandre M. Silva Júnior, Fabiano A. Dorça

Universidade Federal de Uberlândia (UFU)
{alexandre.junior, fabianodor}@ufu.br

Introdução: Um Sistema Tutor Inteligente (STI) é um sistema que incorpora Inteligência Artificial para tutorar o estudante no processo de ensino e aprendizagem. Esse tipo de sistema permite maior grau de individualização, pois adapta suas instruções de acordo com as metas e o nível de entendimento do aluno, além de resolver problemas e explicar como os fez. Os STIs podem ser aplicados para desenvolver as habilidades de pensamento computacional, que é capacidade de raciocinar de maneira lógica e sistemática para resolver problemas e desafios nas diversas áreas do conhecimento, e que se baseia em quatro pilares principais: decomposição, abstração, reconhecimento de padrão e algoritmo. **Objetivos:** Este trabalho propõe o desenvolvimento de um Sistema Tutor Inteligente voltado ao treinamento de habilidades de pensamento computacional. O sistema será capaz de gerar atividades personalizadas com base nos quatro pilares do pensamento computacional, oferecendo instruções e feedbacks personalizados, por meio de uma IA generativa e incorporará elementos gamificados para prover maior engajamento. Ademais, busca-se aprimorar o suporte à aprendizagem e elevar o interesse dos estudantes em disciplinas com altos índices de retenção e dificuldade. **Método:** O método proposto inicia-se focado nas dificuldades dos alunos no aprendizado de computação, a fim de determinar os problemas e assim traçar os objetivos do estudo. Para embasar a solução, realiza-se uma revisão sistemática da literatura sobre STIs, gamificação e abordagens pedagógicas. Com base na revisão, bem como entrevistas e questionários aplicados a professores e alunos, são levantados os requisitos funcionais e não funcionais do sistema. Após coleta e análise dos dados, inicia-se a modelagem da arquitetura e desenvolvimento do STI utilizando uma abordagem incremental e iterativa. Por fim, será feito o processo de validação técnica e pedagógica, sendo esta realizada por meio de estudos com usuários, para aferir o engajamento e eficácia do sistema proposto. **Resultados:** É esperado que esse trabalho contribua para o desenvolvimento das habilidades de pensamento computacional, promovendo um aprendizado mais ativo e personalizado. Além disso, o sistema proposto também deverá favorecer a personalização do ensino e a melhoria no desempenho acadêmico dos estudantes. Contudo, necessita-se de uma rigorosa validação técnica e pedagógica para que se tenha qualidade nos conteúdos gerados por IA. **Conclusão:** O trabalho, embora ainda em desenvolvimento, apresenta a proposta de um Sistema Tutor Inteligente (STI) voltado ao desenvolvimento de habilidade de pensamento computacional. A solução utiliza IA generativa e elementos de gamificação para oferecer um ambiente personalizado e iterativo. Por meio da revisão da literatura e levantamento de dados, busca-se fundamentar o desenvolvimento do sistema, que visa aprimorar o suporte a aprendizagem e estimular o raciocínio lógico dos estudantes.

Trilha: Trabalho de Graduação.

Uma Abordagem Baseada em Transformer para Classificação de Espectros FTIR no Diagnóstico de Câncer Oral

Lucas Procópio, Paulo Souza, Murillo Carneiro, Robinson Silva

Universidade Federal de Uberlândia (UFU)

{lucas.janot, pdodonto, mgcarneiro, robinsonsabino}@ufu.br

Introdução: A espectroscopia no infravermelho por transformada de Fourier se destaca como uma técnica promissora para o diagnóstico não invasivo de doenças, incluindo o câncer oral, cuja alta taxa de mortalidade está associada ao diagnóstico tardio. A análise de espectros biológicos, como os obtidos da saliva, permite identificar alterações associadas à carcinogênese. Contudo, desafios como variabilidade espectral, ruído e pequeno tamanho amostral limitam o desempenho de métodos tradicionais de aprendizado de máquina. Nesse contexto, as arquiteturas Transformer surgem como alternativas promissoras, por sua capacidade de capturar padrões complexos nos dados espectrais. **Objetivos:** Propor e avaliar uma nova arquitetura baseada em Transformer para classificar espectros FTIR de saliva no diagnóstico de câncer oral. Especificamente, o estudo busca: (i) adaptar a arquitetura Transformer ao domínio espectral, (ii) comparar seu desempenho com métodos consolidados e estado da arte como SVM, TabPFN2, CatBoost e XGBoost; e (iii) analisar a interpretabilidade do modelo por meio de mapas de atenção. **Método:** O conjunto de dados é composto por 65 espectros FTIR de amostras salivares obtidas de pacientes com câncer oral (39) e indivíduos saudáveis (26). O pipeline inclui correção da linha de base (Polynomial, Rubberband e ASLS), normalização pela banda Amida I e seleção da região espectral entre 3050–850 cm^{-1} . A validação foi realizada com k-fold estratificado ($k=10$), e as métricas de avaliação: acurácia, precisão, sensibilidade, especificidade e F1-score. **Resultados:** Os experimentos foram realizados com os modelos XGBoost, SVM, TabPFN2 e CatBoost, aplicados a diferentes estratégias de correção de linha de base e normalização. Entre os métodos avaliados, destacou-se a correção por ASLS e a normalização pela Amida I. O XGBoost obteve acurácia de $0,70 \pm 0,18$ e F1 de $0,76 \pm 0,14$, enquanto o SVM alcançou $0,65 \pm 0,14$ e $0,75 \pm 0,10$, respectivamente. O modelo TabPFN2 apresentou resultados consistentes, com acurácia de $0,61 \pm 0,09$ e $0,73 \pm 0,08$. O CatBoost obteve $0,72 \pm 0,17$ e $0,77 \pm 0,14$, respectivamente; a análise de importância de atributos indicou que as regiões lipídica e de fingerprint foram mais discriminativas. Embora o modelo Transformer ainda esteja em desenvolvimento, os resultados confirmam a adequação das etapas de pré-processamento e validação adotadas, estabelecendo uma base para a comparação entre as abordagens e o avanço da arquitetura proposta. **Conclusão:** O estudo estabelece uma base metodológica sólida para a implementação e avaliação de arquiteturas Transformer no diagnóstico não invasivo de câncer oral. A combinação de pré-processamento adequado e aprendizado de atenção espectral pode contribuir para resultados generalizáveis, contribuindo com um grande desafio de estudos na área: a falta de consenso entre metodologias. Como perspectivas futuras, propõe-se ampliar o conjunto de dados e explorar a aplicabilidade do modelo em outros dados biomédicos.

Trilha: Trabalho de Graduação.

Uma Ferramenta Computacional para Geração Automatizada de Perfis Geológicos a partir de Dados KML

Augusto Barbosa, Daniel Caetano, Fabiano Mota

Universidade Federal de Uberlândia (UFU)

{augusto.barbosa, daniel.caetano, fabianomota}@ufu.br

Introdução: A elaboração de perfis geológicos é uma tarefa central em geociências e engenharias, indispensável para o planejamento de infraestrutura e prospecção mineral. Historicamente, este é um processo predominantemente manual, complexo e demorado. Profissionais precisam integrar manualmente dados de mapas topográficos e geológicos, um fluxo de trabalho que consome horas e é suscetível a erros de transcrição e interpolação. Softwares SIG, embora poderosos, apresentam uma curva de aprendizado acentuada e um fluxo de trabalho fragmentado. Ferramentas como o Google Earth Pro extraem perfis de elevação, mas falham em integrar diretamente as camadas geológicas, forçando o pós-processamento manual. **Objetivos:** O objetivo principal deste trabalho foi preencher essa lacuna de eficiência, projetando uma solução de software para automatizar a geração de perfis geológicos. O resultado é o Earthcake, uma aplicação desktop de código aberto, multiplataforma e multilíngue, concebida para ser intuitiva e reduzir a carga cognitiva. O sistema visa permitir que o usuário, a partir de arquivos KML de linha (trajeto) e polígonos (formações geológicas), obtenha automaticamente um perfil altimétrico coeso e colorido (PNG) e os dados brutos (JSON). **Método:** O Earthcake foi desenvolvido em Python, seguindo uma metodologia de prototipagem. A arquitetura utiliza PyQt5 para a interface gráfica (GUI), Matplotlib para a visualização dos gráficos e Shapely para as operações geométricas de ponto-em-polígono. O pipeline de processamento inicia lendo os KMLs. O trajeto da linha é densificado por interpolação linear (usando Geopy) em intervalos de 50 metros. Em seguida, a elevação de cada ponto é obtida via requisições HTTP à API pública Open-Elevation. A etapa de classificação usa Shapely para determinar qual polígono geológico contém cada ponto. Por fim, para produzir um perfil realista, um filtro de média móvel (com NumPy) é aplicado aos dados de elevação, gerando uma curva suavizada. **Resultados:** A validação da ferramenta confirmou que a metodologia de interpolação controlada e suavização produz o perfil de mais alta fidelidade visual no menor tempo de execução. Estudos de caso em Mogi das Cruzes (SP) e Serra Negra (Patrocínio-MG) demonstraram a eficácia e escalabilidade da ferramenta. O perfil de 17 km da estrutura vulcânica de Serra Negra foi gerado em 297 segundos, correlacionando com sucesso a morfologia do relevo com as unidades litológicas. A análise quantitativa por sobreposição demonstrou alta aderência dos dados da Earthcake com os pontos de controle do Google Earth Pro. **Conclusão:** O Earthcake cumpre seu objetivo principal, estabelecendo-se como uma solução de software livre eficaz e direcionada. Ela agiliza e democratiza a geração de perfis geológicos, transformando um processo manual de horas em uma tarefa automatizada de minutos. O projeto aumenta a produtividade de geocientistas e garante maior padronização e reprodutibilidade nas análises.

Trilha: Trabalho de Graduação.

Uma investigação do Algoritmo Genético aplicado ao Sudoku clássico 9×9

Andressa Bernardes, Christiane Brasil

Universidade Federal de Uberlândia (UFU)
{andressa.bernardes, christiane}@ufu.br

Introdução: Os Algoritmos Evolutivos (AEs) são técnicas computacionais inspiradas na Teoria da Evolução de Darwin, amplamente utilizadas em problemas complexos. Seu princípio fundamental consiste em evoluir uma população de soluções candidatas por meio de um processo iterativo, em que são selecionados os indivíduos pais a fim de gerar filhos por meio de operadores genéticos, como recombinação e mutação. Uma das classes mais notáveis de AEs são os Algoritmos Genéticos (AGs), que podem ser utilizados para a otimização de diversos problemas combinatórios. Eles se destacam pelo uso obrigatório do operador de recombinação (crossover) para gerar novas soluções. O quebra-cabeça lógico Sudoku é um forte exemplo de problema combinatório. Embora as regras (sem repetição de números em linhas, colunas e submatrizes 3×3) do Sudoku clássico 9×9 sejam simples, ele é classificado como um problema NP-completo, o que torna o AG uma abordagem promissora para a sua resolução. **Objetivos:** Desenvolver e otimizar um AG para a resolução do jogo Sudoku clássico, investigando como diferentes configurações de parâmetros e a escolha de operadores genéticos impactam seu desempenho. **Método:** Foi implementado um AG em Python, onde cada indivíduo é uma matriz 9×9. A função de fitness foi definida como a contagem total de conflitos (números repetidos em linhas, colunas ou submatrizes), objetivando minimizar esse valor a zero. A geração inicial foi implementada puramente aleatória, sorteando os números para as células vazias. A metodologia envolveu experimentação controlada para calibrar parâmetros e selecionar os melhores operadores, comparando diferentes abordagens para seleção (torneio, roleta), recombinação (uniforme, ponto único) e mutação (uniforme, por troca, por gene), além da aplicação de elitismo. **Resultados:** A calibração experimental definiu a configuração de melhor eficácia: população de 600 indivíduos, taxa de mutação de 1%, taxa de elitismo de 5% e seleção por torneio com tamanho 3. A análise dos operadores foi crucial, e a combinação de Recombinação Uniforme e Mutação por Gene apresentou o melhor desempenho. Em uma comparação com o trabalho de referência para o mesmo problema — cujos métodos incluíam mutação por troca, recombinação de ponto único, seleção por roleta e geração inicial aleatória sem repetição nas submatrizes — o modelo proposto, com sua configuração otimizada, mostrou resultados superiores em relação ao fitness. Ele alcançou taxas de sucesso de até 30% em instâncias fáceis, enquanto a réplica do trabalho de referência não obteve nenhuma solução nas mesmas condições. **Conclusão:** O estudo determinou uma configuração otimizada para o AG, provando que o ajuste metodológico de parâmetros e a escolha de operadores são os fatores mais determinantes para o desempenho. A configuração final validou a eficácia do AG para instâncias de nível fácil, superando a abordagem de referência.

Trilha: Trabalho de Graduação.

Uma solução baseada em Modelos de Linguagem e Geração Aumentada de Recuperação para auxílio ao discente da disciplina de Sistemas Operacionais

Fabíola Pereira, Ana Julia Simoes

Universidade Federal de Uberlândia (UFU)
{fabiola.pereira, ana.simoess1}@ufu.br

Introdução: O trabalho que está a ser desenvolvido é uma solução que se baseia em uma assistente virtual, uma ferramenta que dará um auxílio extra aos discentes e também aos docentes no ensino da disciplina de Sistemas Operacionais (SO) do curso de Bacharelado em Sistemas de Informação. Foi feita uma pesquisa entre as matérias do curso de Sistemas de Informação para verificar qual seria melhor aplicada a esta pesquisa, e o resultado foi a disciplina de Sistemas Operacionais, pela quantidade de conceitos teóricos computacionais que necessitam de um entendimento mais aprofundado sobre o funcionamento dos sistemas computacionais. Também foi feita uma pesquisa sobre ferramentas similares a esta, mas o resultado foi que nenhuma é específica para a disciplina, e a maioria apenas faz extração de texto, o que gera respostas com certo tipo de limitação e também alucinações. Para esta pesquisa, utiliza-se Retrieval-Augmented Generation (RAG) para reduzi-las. **Objetivos:** O objetivo esperado nesta pesquisa é desenvolver uma ferramenta eficiente que melhore o entendimento da disciplina, que possa ser usada pelo discente para tirar dúvidas específicas e realizar trabalhos com um conteúdo mais centralizado na matéria. A justificativa central está na necessidade de uma solução acessível e de fácil uso para a faculdade, que tanto o discente quanto o docente possam utilizar de forma eficaz. **Método:** Será utilizado o método de Modelos de Linguagem de Larga Escala (LLMs) com Retrieval-Augmented Generation (RAG), implementado em Python e com interface desenvolvida no framework React. Com essas ferramentas, espera-se realizar uma recuperação de informação a partir de diversas fontes, não apenas as que o docente pode disponibilizar, mas também fontes que podem ser encontradas na internet e em outros repositórios. A interface desenvolvida em React será intuitiva e de fácil entendimento. Após a conclusão da ferramenta, será feito um experimento com os discentes, que utilizarão a ferramenta em casos reais de dúvidas e estudos. Com isso, será analisado se ela é eficaz e se os alunos se adaptaram. **Resultados:** Ainda assim, os resultados previstos indicam que a adoção dessa ferramenta promoverá maior engajamento, autonomia e eficiência no aprendizado. **Conclusão:** O que se espera é um estudo mais eficaz, com respostas adaptadas para que o discente obtenha respostas precisas. Espera-se demonstrar que o uso de modelos de linguagem LLMs e RAG oferece uma abordagem superior às abordagens tradicionais de extração textual, proporcionando maior precisão e contextualização nas respostas.

Trilha: Trabalho de Graduação.

Uso de Heurísticas na Detecção de Mixer na Blockchain Bitcoin

Maria Rita V. Souza, Ivan S. Sendin

Universidade Federal de Uberlândia (UFU)

{mariarvsouza, sendin}@ufu.br

Introdução: O Bitcoin é uma moeda digital descentralizada que utiliza a tecnologia peer-to-peer (P2P) para operar sem uma autoridade central, garantindo transparência, pseudoanonimato e acessibilidade. Essas características, aliadas à forte criptografia e relativa anonimidade conferem ao Bitcoin uma confiabilidade que pode, paradoxalmente, ser usada como aliada para atividades ilícitas. Porém, o pseudoanonimato dessa criptomoeda garante que, embora os endereços não possuam as identidades dos usuários explícita e propriamente ditas, as transações e seus dados estão registrados na blockchain, o que permite que inúmeras análises sejam feitas sobre as transações e seus envolvidos, mesmo quando adotam técnicas avançadas para aumentar seu anonimato, como os mixers, que são serviços que recebem e misturam Bitcoins de múltiplos usuários para quebrar o vínculo entre os endereços de origem e os de destino e dificultar o rastreamento de recursos ilícitos. Neste trabalho, apresentamos o resultado da aplicação de heurísticas na detecção de dois serviços populares no Bitcoin no primeiro semestre de 2025, o JoinMarket e o Wasabi. **Objetivos:** O objetivo principal desta pesquisa é investigar a aplicação de heurísticas na detecção de serviços de mixers no Bitcoin, para aprimorar as técnicas de análise forense. Isso inclui analisar o funcionamento dessa criptomoeda, explorar o uso de mixers como ferramenta para cashing out, e identificar e caracterizar transações de dois mixers existentes. **Método:** Um estudo bibliográfico sobre os mixers Wasabi e JoinMarket permitiu detectar características de transações executadas por estes serviços e desenvolver heurísticas para a sua identificação. **Resultados:** Foram coletadas um total de 69.366.786 transações, referentes ao primeiro semestre de 2025, e, entre elas, foram extraídas no total 8.611 transações de JoinMarket e 23.343 transações de Wasabi no período analisado. A partir desses dados, foi feita uma análise sobre a quantidade de ocorrências de cada mixer por dia nesse intervalo de tempo, o que resultou em uma média de aproximadamente 128,97 transações diárias de Wasabi e 47,57 de JoinMarket. **Conclusão:** Os resultados iniciais indicam que as heurísticas selecionaram corretamente as transações de mixer. As próximas etapas do trabalho são a validação das heurísticas usando base de dados existentes na literatura e a análise das transações.

Trilha: Trabalho de Graduação.

Utilização de Redes Neurais com Otimização de Processamento de Imagens para detecção de Deficiências Nutricionais em Folhas de Café

Luanna Rabelo, Thiago Ribeiro

Universidade Federal de Uberlândia (UFU)

{luannarabelo, tpribeiro}@ufu.br

Introdução: O café chegou ao Brasil em 1727 e rapidamente se espalhou entre os agricultores, fazendo do país o maior produtor de café do mundo desde 1820 até hoje. Entretanto, com o consumo sempre crescente e as constantes mudanças climáticas, além das alterações no solo e no ambiente, o cultivo do café passou a enfrentar diferentes tipos de deficiências nutricionais e doenças que afetam suas folhas e produção. Diante desse cenário, a aplicação de técnicas avançadas de processamento digital de imagens, aliadas a redes neurais artificiais, surge como uma solução eficiente e necessária para a detecção dessas deformidades nas plantações de café. **Objetivos:** O principal objetivo deste estudo é aprimorar a detecção de deficiências nutricionais em folhas de café, empregando técnicas avançadas de processamento de imagens e aprendizado de máquina. Com o intuito de prevenir danos às plantações de café decorrentes de baixa nutrição por meio da análise de suas folhas, e assim contribuir para o desenvolvimento de estratégias eficazes no manejo e cuidado das culturas. **Método:** Foram analisadas 669 imagens de folhas cultivadas em ambiente controlado com deficiências nutricionais selecionadas: 96 de cálcio (Ca), 76 de potássio (K), 95 de magnésio (Mg), 59 de nitrogênio (N), 102 de fósforo (P) e 99 de enxofre (S). Além dessas, foram trabalhadas 142 imagens sem nenhuma deficiência nutricional. As imagens passaram por um processo de filtragem, visando realçar e remover características relevantes, seguido pela segmentação, que consiste na divisão da imagem em regiões distintas para facilitar a análise dos elementos. Em seguida, foram aplicados filtros e padrões para destacar, ocultar e modificar áreas específicas das fotos originais, bem como técnicas de extração de características, permitindo a identificação e o isolamento de padrões relevantes para aumentar a precisão dos modelos de aprendizado. Por fim, as imagens foram agrupadas e rotuladas conforme o tipo de deficiência nutricional, permitindo o treinamento supervisionado das redes neurais. O processamento foi realizado utilizando Redes Neurais Artificiais (RNA) e abordagens de Deep Learning, especialmente Redes Neurais Convolucionais (CNNs) e a arquitetura Xception, adequadas para o reconhecimento de padrões em imagens. Também foi testada a arquitetura Vision Transformer (ViT), que aplica mecanismos de atenção para identificar relações espaciais entre as partes da imagem. **Resultados:** A aplicação das técnicas descritas junto à arquitetura Xception resultou em uma acurácia de aproximadamente 64,94%, desempenho muito próximo ao obtido com a técnica Vision Transformer (ViT), que alcançou 65%. **Conclusão:** Ambas as técnicas de aprendizado de máquina apresentaram resultados satisfatórios. No entanto, o desempenho poderia ser aprimorado com o uso de um banco de imagens maior e mais detalhado, permitindo um treinamento mais robusto e preciso dos modelos.

Trilha: Trabalho de Graduação.

Trilha: Trabalhos de Pós-graduação

Algoritmos Evolutivos para a Otimização Dinâmica de um Problema Discreto com Muitos Objetivos

Thiago Lafetá, Luiz Gustavo Martins

Universidade Federal de Uberlândia (UFU)

{thiago.fialho, lgamartins}@ufu.br

Introdução: Diversos problemas de otimização presentes no mundo real apresentam natureza dinâmica e envolvem múltiplos objetivos conflitantes. Embora algoritmos evolutivos tenham sido amplamente aplicados a problemas estáticos, a combinação simultânea de dinamismo e múltiplos objetivos (DMOPs) ainda representa um desafio significativo para a área. **Objetivos:** O objetivo deste trabalho é desenvolver e avaliar novos algoritmos evolutivos multiobjetivo dinâmicos (DMOEAs) capazes de lidar de forma eficiente com ambientes em mudança e múltiplos objetivos conflitantes. Para isso, são propostos três algoritmos: D-MEANDS, D-MEANDS-MD e D-MEANDS-II, todos inspirados nos algoritmos MEANDS e MEANDS-II originalmente aplicados a problemas estáticos e discretos. **Método:** As versões dinâmicas propostas incorporam estratégias de memória, mecanismos de preservação de diversidade e aprimoramentos no gerenciamento de subpopulações para lidar com mudanças ambientais. A avaliação experimental dos algoritmos foi realizada utilizando o Dynamic Multiobjective Knapsack Problem (DMKP), considerando instâncias com até oito objetivos e cenários contendo vinte mudanças de ambiente. **Resultados:** Os experimentos conduzidos demonstraram que os algoritmos propostos apresentam desempenho competitivo em relação aos métodos da literatura. Em diversos cenários, os DMOEAs desenvolvidos superaram abordagens já consolidadas, evidenciando sua eficácia na adaptação a ambientes dinâmicos e multiobjetivo. **Conclusão:** A incorporação de estratégias dinâmicas aos algoritmos MEANDS e MEANDS-II resultou em versões capazes de lidar de maneira eficaz com DMOPs. Os resultados obtidos demonstram o potencial dos algoritmos propostos, destacando-os como alternativas promissoras no avanço da otimização evolutiva para DMOPs.

Trilha: Trabalho de Pós-graduação.

Análise e Detecção de Anomalias de Cobertura em Redes Celulares: Uma Abordagem com Aumento de Dados e Modelagem Preditiva

Daniel Ricardo Oliveira¹, Flávio de Oliveira Silva²

¹Universidade Federal de Uberlândia (UFU)

²Universidade do Minho (UMinho)

drcoliveira@ufu.br, flavio@di.uminho.pt

Introdução: Com a evolução das redes móveis, agregação de múltiplas faixas de frequências e necessidade de interoperabilidade entre sistemas, o processo de ajuste e otimização de cobertura celular se tornou extremamente complexo. Tradicionalmente, o método utilizado é o da realização de drive tests (medições georreferenciadas com equipamentos em carros) e posterior análise por equipes de especialistas. Neste trabalho, propõe-se o ARCADE (Automated Radio Coverage Anomalies Detection and Evaluation), uma metodologia de otimização auto-contida do sistema celular, na qual a própria rede consegue diagnosticar falhas de cobertura ou interferência sem intervenção humana ou integração com sistemas externos à rede. **Objetivos:** O objetivo geral da tese é a de propor uma metodologia para identificação de anomalias de cobertura celular de forma automatizada, contando apenas com dados georreferenciados de cobertura por célula, sem o auxílio de informações externas. Isso implica na não utilização de modelos empíricos de predição matemática de cobertura, de forma que o modelo deve se basear apenas em dados reais. Além disso, os dados de cobertura podem ser esparsos e não abranger totalmente a área analisada, ou seja, há necessidade de extrapolação desses dados. **Método:** A metodologia ARCADE divide-se em módulos que vão da aquisição dos dados até a identificação final das anomalias de cobertura, passando por pré-processamento dos dados, aumento de amostras através de processos gaussianos com kernel espacial e treinamento e generalização do modelo com o uso de redes neurais artificiais. O ARCADE se diferencia das atuais propostas principalmente por não contar com entrada os dados do projeto da rede, geográficos ou indicadores de desempenho e não se apoiar em modelos analíticos (predições matemáticas teóricas), baseando suas inferências somente em dados físicos. Essas restrições permitem que a rede móvel por si só seja capaz de autodiagnosticar-se, de modo que o ARCADE confere ao sistema uma autopercepção sobre o próprio desempenho de cobertura. **Resultados:** Uma implementação prática da metodologia levou a resultados muito expressivos. De dados coletados de uma rede móvel operacional, contendo amostras esparsas, anômalas e outliers, o ARCADE conseguiu, com sucesso, gerar mapas de cobertura que representam muito fielmente uma condição real da cobertura em campo. Além disso, cruzando-se as informações extrapoladas, foi possível identificar quais células apresentaram comportamento anormal, sem que houvesse intervenção humana especializada na análise. **Conclusão:** A proposta da metodologia ARCADE busca idealmente uma autopercepção da cobertura celular pela própria rede, o que é um desafio devido à baixa disponibilidade de dados, amostras esparsas, inerente sobreposição e entrelaçamento das células e consequente complexidade dessa análise. As primeiras avaliações mostraram-se muito promissoras, já com resultados concretos e aplicáveis no contexto da indústria.

Trilha: Trabalho de Pós-graduação.

Arquitetura Híbrida para Detecção fim a fim de Exfiltração DNS com Inteligência em Camada de Rede e Inspeção em Nuvem para Análise do Tráfego Criptografado

Cristiano Borges, Rodrigo Miani

Universidade Federal de Uberlândia (UFU)

cborges@fortinet.com, miani@ufu.br

Introdução: O aumento dos ataques de exfiltração de dados utilizando o protocolo Domain Name System (DNS) representa um dos maiores desafios na segurança cibernética moderna. O tunelamento de DNS, especialmente em sua versão criptografada DNS over HTTPS (DoH), permite que agentes maliciosos transmitam informações sensíveis por meio de consultas aparentemente legítimas, contornando mecanismos tradicionais de proteção. Esse cenário torna a detecção de exfiltração via DNS uma tarefa complexa, uma vez que o tráfego é considerado essencial para o funcionamento da rede e, portanto, raramente inspecionado de forma profunda. **Objetivos:** Este trabalho propõe o desenvolvimento de uma arquitetura híbrida denominada Arsenal-DLP, voltada à detecção fim a fim de exfiltração de dados via DNS. A solução integra inspeção em tempo real na camada de rede com análise avançada em nuvem, combinando aprendizado de máquina (Machine Learning, ML) e técnicas de Explainable Artificial Intelligence (XAI). O objetivo é oferecer uma abordagem escalável e transparente capaz de identificar padrões anômalos em fluxos DNS, inclusive em tráfego criptografado, superando limitações observadas em soluções tradicionais baseadas em assinaturas. **Método:** O método proposto adota uma abordagem aplicada e exploratória. A arquitetura é composta por duas camadas interconectadas: a camada local, responsável pela captura e análise em tempo real de pacotes DNS e fluxos suspeitos, e a camada em nuvem, dedicada à análise aprofundada, descriptografia e aplicação de modelos de ML. Um ciclo de feedback contínuo conecta ambas as camadas, permitindo que as descobertas na nuvem atualizem dinamicamente as políticas de segurança locais. Além disso, o trabalho se propõe a criar um conjunto de dados inédito, resultante da junção de cenários simulados de tráfego DNS tradicional (Do53) e criptografado (DoH), contribuindo para preencher uma lacuna existente na literatura sobre detecção de exfiltração de dados em DNS. **Resultados:** Os resultados esperados incluem a criação de um modelo capaz de identificar exfiltrações com alta precisão, mesmo em tráfego criptografado, e a redução de falsos positivos. A explicabilidade fornecida pelo XAI aumenta a confiança dos analistas e a auditabilidade das decisões do sistema. **Conclusão:** Em comparação com os trabalhos presentes na literatura, a proposta se diferencia pela integração entre camadas de análise, pelo uso de aprendizado explicável e pela construção de um dataset híbrido que amplia a representatividade dos cenários de ataque, consolidando uma solução prática e inovadora para mitigar vazamentos de dados via DNS.

Trilha: Trabalho de Pós-graduação.

Avaliação de Fertilidade Seminal: Uma Abordagem Centrada na Análise dos Parâmetros de Morfologia, Vitalidade e Contagem Espermáticas com Base em Imagens Digitais e Deep Convolutional Neural Networks

Saide Saide, Marcelo Beletti, Bruno Travençolo

Universidade Federal de Uberlândia (UFU)

{saide.saide, mebeletti, travencolo}@ufu.br

Introdução: A fertilidade masculina é um processo biológico complexo definido pela capacidade dos espermatozoides fecundarem o óvulo feminino. No contexto da indústria alimentícia, a reprodução constitui um fator determinante para a manutenção da produção da carne e seus derivados. Tradicionalmente, a análise da fertilidade masculina é realizada manualmente por profissionais especializados via observação laboratorial de amostras de sêmen. Entretanto, a avaliação manual apresenta limitações relacionadas à subjetividade e à variabilidade dos resultados. Como contramedida, foram introduzidas soluções de análise computadorizadas tais como sistemas CASA, técnicas de processamento digital de imagens e recentemente as abordagens de DCNN. Porém, os sistemas CASA são extremamente caros, as técnicas de processamento digital de imagens deparam-se com limitações como baixa generalização, enquanto as soluções baseadas em DCNN analisam majoritariamente fertilidade masculina humana. **Objetivos:** O presente trabalho tem como principal objetivo desenvolver e validar um método computacional de avaliação da fertilidade seminal bovina, baseado em algoritmos de DCNNs aplicados à análise automatizada de imagens digitais de sêmen. **Método:** A pesquisa integra três etapas: estudo, implementação, e experimentos, que serão desenvolvidas num processo iterativo. A primeira etapa (estudo) compreende a exploração, análise e compreensão das imagens (seminais) disponíveis, com vista a construção do conjunto de dados para alimentação dos algoritmos na fase de treinamento, validação e teste, assim como, a exploração, análise e compressão dos diversos algoritmos de adequados à realização das tarefas de visão computacional. A segunda etapa (implementação), consiste na programação do método mencionado na presente proposta, resultantes das descobertas da etapa anterior (estudo). A terceira etapa (experimentos), concentra-se no treinamento, validação e teste dos algoritmos de DCNNs, para obtenção dos resultados esperados. **Resultados:** Os principais resultados parciais obtidos incluem: um conjunto criado de 4116 imagens e igual número de máscaras para segmentação, treinamento, validação e teste de uma U-Net padrão cujo desempenho atingiu uma média de 0,95 de coeficiente de dice de treino, validação e teste, e uma média de 0,05 de loss de treino, validação e teste. **Conclusão:** Os resultados preliminares demonstram que as DCNNs podem superar as limitações das abordagens atuais de análise de fertilidade masculina. Espera-se que os resultados definitivos contribuam para o avanço do estado da arte na análise da fertilidade masculina e da visão computacional, assim como, fortaleçam o setor pecuário, promovendo maior eficiência reprodutiva, sustentabilidade produtiva e retorno econômico.

Trilha: Trabalho de Pós-graduação.

Avaliação do Impacto de Metodologias Ativas e Pensamento Computacional no Desempenho em Programação Competitiva

Camila Santos, Rafael Araújo, João Pereira

Universidade Federal de Uberlândia (UFU)

{camilacruz, rafael.araujo, joaohs}@ufu.br

Introdução: O aprimoramento do ensino de Computação é essencial para o desenvolvimento de competências digitais e cognitivas alinhadas à BNCC e à Educação 5.0. Nesse contexto, a Programação Competitiva, exemplificada pela Olimpíada Brasileira de Informática (OBI), destaca-se como uma estratégia para estimular o Pensamento Computacional (PC), o raciocínio lógico e a resolução criativa de problemas. O presente estudo aborda o desafio de desenvolver e validar um método de treinamento estruturado e contínuo no Instituto Federal do Triângulo Mineiro (IFTM), buscando superar o baixo desempenho institucional histórico na OBI. **Objetivos:** O objetivo central é avaliar o impacto dessa integração sobre o desempenho dos estudantes do IFTM na OBI. Especificamente, busca-se validar a eficácia do modelo pedagógico desenvolvido ao longo de três ciclos (2023–2025), promover o aperfeiçoamento contínuo do método conforme as necessidades dos alunos iniciantes e ampliar a participação institucional. **Método:** A pesquisa adota uma abordagem híbrida e exploratória, baseada no Estudo de Caso e na Pesquisa Baseada em Design (DBR), caracterizada por ciclos iterativos de aplicação, análise e refinamento. Essa estrutura possibilitou o aprimoramento progressivo do método em contexto real, unindo teoria e prática de modo colaborativo. O treinamento ocorreu em formato híbrido (presencial e remoto), com trilhas de estudo em C++ e o uso de plataformas como a Neps Academy, centradas na resolução de problemas da OBI. **Resultados:** A avaliação do impacto considerou os resultados oficiais da OBI e análises estatísticas das notas e classificações dos participantes. Os resultados confirmam avanços estatisticamente significativos em 2025: o número de participantes cresceu 60,61%, passando de 66 em 2023 para 106 em 2025, enquanto o número de classificados para a segunda fase triplicou, de 20 para 60 (aumento de 200%). A média das pontuações evoluiu de 111,46 para 187,64 ($p < 0,0001$), demonstrando melhoria no desempenho. A participação feminina também aumentou, de uma estudante classificada em 2023 para nove em 2025. Apesar da redução no número de campi participantes, devido à greve institucional e a desafios logísticos, os resultados evidenciam a consolidação do método, especialmente no Campus Uberlândia Centro. **Conclusão:** A pesquisa confirma que a Programação Competitiva, aliada ao PC e à PBL, é uma estratégia pedagógica eficaz para o ensino de Computação, resultando em ganhos significativos de desempenho e engajamento. A principal contribuição metodológica está na aplicação da DBR, que promoveu a evolução contínua do modelo e a integração de ferramentas tecnológicas que fortalecem a autonomia e a metacognição. Os resultados de 2025 consolidam a eficácia e a maturidade do método, demonstrando seu potencial de replicação e sua relevância como modelo de referência nacional para o ensino de Computação mediado por competições científicas.

Trilha: Trabalho de Pós-graduação.

Avaliações de Modelos de Deep Learning na Classificação Temporal de Deficiências Nutricionais em Folhas de Café

Hericles Ferraz, Maurício Escarpinati, Thiago Ribeiro

Universidade Federal de Uberlândia (UFU)

{hericles.ferraz, mauricio, tpribeiro}@ufu.br

Introdução: A cafeicultura, pilar da agricultura brasileira, enfrenta desafios agrônômicos que ameaçam sua produtividade, como as deficiências nutricionais. O diagnóstico visual tradicional é subjetivo, lento e muitas vezes tardio, ocorrendo quando o potencial produtivo foi comprometido. A agricultura de precisão, aliada a deep learning, surge como alternativa promissora para automatizar e antecipar a detecção. No entanto, a maioria dos estudos trata o diagnóstico como uma classificação estática, usando imagens de sintomas avançados e ignorando a evolução temporal dos sinais. Esta lacuna motivou o desenvolvimento de um framework para análise temporal, visando identificar a partir de que momento os sinais se tornam discriminativos e qual arquitetura é mais adequada para a detecção precoce.

Objetivos: O objetivo geral é comparar sistematicamente o desempenho de arquiteturas de deep learning, Redes Neurais Convolucionais (CNNs), Vision Transformers (ViTs) e híbridos CNN-Transformer, na classificação temporal de múltiplas deficiências nutricionais (N, P, K, Ca, Mg, S) em folhas de café, a partir de imagens RGB. Como objetivos específicos, busca-se determinar a janela temporal em que os sintomas se tornam discriminativos; avaliar se arquiteturas baseadas em atenção (ViT e híbridas) superam as convolucionais em diferentes fases do desenvolvimento sintomático e validar se os modelos focam em regiões foliares agronomicamente relevantes via técnicas de interpretabilidade (Grad-CAM). **Método:** Utilizou-se um dataset original, composto por imagens de folhas de café coletadas ao longo de cinco meses sob condições controladas, representando sete classes (seis deficiências e um controle). O dataset foi particionado em três fases temporais: Inicial (primeiros 2 meses), Intermediária (mês 3 e 4) e Tardia (último mês). Foram treinados e avaliados modelos CNN (ResNet50), ViT (ViT-B/16) e uma arquitetura híbrida, utilizando métricas como acurácia, F1-Score e matriz de confusão. **Resultados:** O desempenho dos modelos mostrou-se fortemente dependente do estágio dos sintomas, com os melhores resultados ocorrendo na fase tardia. A arquitetura puramente convolucional ResNet50 superou os modelos híbridos com 77,1% de acurácia em uma resolução de 512x512 pixels. Esses resultados sugerem que a extração de características locais (CNNs) é mais eficaz que os mecanismos de atenção global para este cenário, e que uma maior resolução tende a fornecer detalhes cruciais para a discriminação. **Conclusão:** Este trabalho propõe uma abordagem metodológica para a análise temporal de deficiências nutricionais, preenchendo uma lacuna ao comparar sistematicamente arquiteturas de ponta em diferentes estágios dos sintomas. Os resultados fornecem insights sobre os limites e a viabilidade da detecção precoce com imagens RGB, contribuindo para o avanço de ferramentas de monitoramento na agricultura de precisão e estabelecendo um protocolo de validação mais rigoroso para estudos futuros na área.

Trilha: Trabalho de Pós-graduação.

Cenários de Aplicação dos Knowledge Graphs em Sistemas de Recomendação: Uma Revisão Sistemática da Literatura

Leandra Vale¹, João Henrique S. Pereira¹, Flávio O. Silva²

¹Universidade Federal de Uberlândia (UFU)

²Universidade do Minho (Uminho)

{leandra.vale, joaohs}@ufu.br, flavio@di.uminho.pt

Introdução: Os sistemas de recomendação (SR) têm ganhado atenção significativa para personalizar e melhorar a experiência dos usuários, especialmente em cenários de sobrecarga de informações. Nos últimos anos, a incorporação de Knowledge Graphs (KGs) aos SR tem atraído considerável interesse, principalmente pela promessa de mitigar alguns problemas críticos dos sistemas tradicionais de recomendação. Como os SR estão presentes em diversos cenários de aplicação, do entretenimento à segurança, à educação e à saúde, torna-se necessário conhecer quais são os principais cenários de aplicação dos KGs em SR. **Objetivos:** Portanto, esta revisão da literatura teve como principal objetivo encontrar resposta para: Quais são os principais cenários de aplicação dos KGs em SR? **Método:** Assim sendo, para avaliar criticamente as pesquisas relevantes nesse contexto, realizou-se uma revisão sistemática da literatura em três etapas: Planejamento, Execução e Análise. Na etapa de planejamento foi escolhido o protocolo PRISMA para o processo de coleta de dados, permitindo que os estudos fossem, na etapa de execução: identificados, selecionados e incluídos, para então serem avaliados e discutidos na etapa de análise. Na etapa de execução, coletou-se os dados do Google Scholar, considerando estudos de aplicação dos KGs em SR publicados até março de 2025. Foram definidos critérios de inclusão e exclusão dos estudos, como a seleção apenas de publicações em periódicos e conferências com acesso disponível ao texto completo em inglês, que não sejam revisões da literatura e abordem cenários de aplicação dos KGs em SR. **Resultados:** Foram identificados inicialmente 140 estudos relevantes, dos quais excluiu-se 6 duplicados, sendo portanto, selecionados 134 estudos. Dentre os selecionados, foram excluídas 13 revisões de literatura, 25 livros, capítulos, teses, preprints ou tutoriais, 2 em outro idioma, 2 que o acesso ao texto completo não estava disponível, 13 nos quais o título e resumo não abordaram a questão de pesquisa. Assim, 79 estudos foram incluídos para análise quanto aos cenários de aplicação dos KGs em SR. **Conclusão:** Após a análise dos estudos selecionados, considerando os critérios estabelecidos, foi possível responder quais os principais cenários de aplicação dos KGs em SR. Dos 79 estudos analisados, 32 eram de propósito geral com aplicação em qualquer cenário prático, 21 envolveram a aplicação de KGs em SR no cenário educacional como a recomendação de cursos, recursos de aprendizagem, caminhos de aprendizagem, vocabulário, pontos de conhecimento, perfil do estudante, tópicos de programação, dentre outros. Os outros 26 estavam distribuídos em diferentes cenários, como aquicultura, carreira, serviços em nuvem, comunicação, cultura e entretenimento, comércio eletrônico, financeiro, saúde, marítimo, programação, psicologia, recrutamento, redes sociais e turismo. Os cenários mais relevantes foram saúde e turismo, seguidos por e-commerce, redes sociais e entretenimento.

Trilha: Trabalho de Pós-graduação.

Classificação de Ameloblastoma e Carcinoma Ameloblástico em Patologia Digital: Uma Abordagem Eficiente com Modelos Híbridos Leves

Domingos L. L. Oliveira¹, Adriano B. Silva²,
Paulo R. Faria², Marcelo Z. Nascimento²

¹Instituto Federal de Educação, Ciência e Tecnologia de São Paulo (IFSP)

²Universidade Federal de Uberlândia (UFU)

domingos.oliveira@ifsp.edu.br

{adriano.barbosa, paulo.faria, marcelo.nascimento}@ufu.br

Introdução: O diagnóstico diferencial entre ameloblastoma (AM) e carcinoma ameloblástico (AC) é um desafio clínico devido à sobreposição morfológica. A digitalização de lâminas inteiras (WSI) na patologia digital permite o uso de sistemas de diagnóstico auxiliado por computador (CADx), mas o tamanho massivo das imagens exige modelos eficientes para análise de patches. Modelos de aprendizado profundo, como redes neurais convolucionais (CNNs), capturam padrões locais, enquanto Vision Transformers (ViTs) modelam o contexto global, mas exigem muitos dados. Modelos híbridos CNN-Transformer surgem como solução, mas há escassez de estudos computacionais para AM e AC devido à raridade dos tumores.

Objetivos: O objetivo foi avaliar o desempenho de três modelos híbridos leves (EdgeNeXt, EfficientViT e LeViT) na classificação de AM e AC em 20× e 40×, analisando a relação entre magnificação, acurácia e o impacto da arquitetura na extração de padrões. **Método:** A partir de 16 casos corados com H&E e anotados por especialistas, foram extraídas 3.744 regiões de interesse (ROIs) em 20× e 14.976 em 40×. Os modelos foram treinados com otimizador Adam e imagens de 224×224 pixels, com divisão por paciente. O desempenho foi avaliado por Área Sob a Curva (AUC), F1-score e Acurácia, e as diferenças comparadas pelo teste de Friedman. **Resultados:** Os três modelos apresentaram melhor desempenho em 40×, confirmando a importância dos detalhes citológicos. O LeViT destacou-se (AUC 0,94, F1 0,84, Acurácia 0,84), seguido do EfficientViT (AUC 0,92, F1 0,83, Acurácia 0,83) e do EdgeNeXt (AUC 0,81, F1 0,69, Acurácia 0,69). Em 20×, os resultados foram: LeViT (AUC 0,83, F1 0,78, Acurácia 0,78), EfficientViT (AUC 0,82, F1 0,72, Acurácia 0,73) e EdgeNeXt (AUC 0,84, F1 0,77, Acurácia 0,78). Notavelmente, o desempenho em 20× do LeViT e EdgeNeXt (Acurácia 0,78) superou o estado da arte anterior (ResNet50 com acurácia de 0,75 e VGG16 com 0,65). Todos os modelos mantiveram custo computacional inferior a 0,5 GMAC e <10M de parâmetros, uma vantagem drástica frente aos 4.1 GMACs (ResNet50) e 15.5 GMACs (VGG16). As diferenças entre os modelos foram estatisticamente significativas ($p < 0,05$) apenas em 40×. **Conclusão:** Modelos híbridos leves podem atingir desempenho competitivo em 40×, onde a atenção global (Transformers) complementa a extração local (CNNs). Os modelos leves superaram o estado da arte (ResNet50, VGG16) para classificação entre AM e AC em acurácia, oferecendo uma redução drástica no custo computacional. A escassez de amostras limita a generalização, reforçando a necessidade de ampliar o dataset e explorar abordagens multiescala.

Trilha: Trabalho de Pós-graduação.

Entre Rimas e Algoritmos: Uma Investigação sobre Tradução Automática Poética

Beatriz Borges, Elaine Paiva, Paulo Gabriel

Universidade Federal de Uberlândia (UFU)

{biarborges, elaine, phrg}@ufu.br

Introdução: A tradução de poesia apresenta desafios específicos por envolver não apenas a transferência de sentido, mas também a preservação de ritmo, rima, estilo e efeitos expressivos. Modelos de tradução automática, sejam baseados em arquiteturas especializadas ou em modelos de linguagem de grande escala, ainda apresentam limitações ao lidar com tais elementos. Embora avanços recentes tenham melhorado a tradução de textos gerais, a tradução poética permanece pouco investigada, especialmente em contextos multilíngues envolvendo línguas latinas. **Objetivos:** O trabalho tem como objetivo comparar o desempenho de diferentes modelos de tradução automática na tradução de poemas, considerando não apenas a similaridade semântica, mas também a preservação temática e estilística. Além disso, buscou-se investigar estratégias de adaptação por meio de fine-tuning com corpora de poemas, músicas e sua combinação, avaliando o impacto dessa adaptação na qualidade da tradução. **Método:** Foi proposto um framework de avaliação em três fases: métricas automáticas (BLEU, METEOR, BERTScore), análise temática com modelagem de tópicos e avaliação humana especializada. Foram analisados modelos especializados (MarianMT, mBART, OpenNMT RNN e Google Translate) e modelos de linguagem (ChatGPT-3.5 e Maritaca AI) em seis pares linguísticos (PT-FR, PT-EN, FR-EN, FR-PT, EN-FR, EN-PT). O método incluiu três módulos experimentais: uso de modelos pré-treinados (Módulo 1) usando o framework proposto, fine-tuning com músicas (Módulo 2) e comparação entre fine-tuning com poemas, com poemas+músicas e sem fine-tuning (Módulo 3). **Resultados:** Os resultados quantitativos revelaram que os modelos de linguagem, juntamente com o Google Translate, superaram os modelos especializados em preservação de sentido e fluência, enquanto o OpenNMT (RNN) apresentou o pior desempenho. A modelagem de tópicos indicou que esses sistemas de maior desempenho mantêm maior consistência temática. A avaliação humana confirmou esses achados ao mostrar que LLMs atingem melhores níveis de sentido e fluidez; entretanto, todos os modelos demonstraram baixo desempenho na preservação de características poéticas, como ritmo, rima e estilo. Nos experimentos de fine-tuning, apenas o mBART apresentou ganhos consistentes com dados específicos, enquanto MarianMT e OpenNMT não se beneficiaram das adaptações. **Conclusão:** O estudo demonstra que as LLMs apresentam melhor desempenho que os modelos especializados de tradução automática na preservação de sentido e fluência em traduções poéticas. Entretanto, a reprodução de elementos estruturais e estilísticos continua sendo uma limitação comum a todos os sistemas avaliados. Avanços futuros dependem do uso de corpora maiores e mais adequados ao domínio poético, tanto para avaliação quanto para fine-tuning, além da ampliação da avaliação humana. Também será necessário investigar estratégias e modelos que incorporem a representação de aspectos formais e estilísticos da poesia.

Trilha: Trabalho de Pós-graduação.

Especialização Supervisionada de Modelos de Linguagem para Decisões Estratégicas no Pôquer

Lucas Diniz, Murillo Carneiro

Universidade Federal de Uberlândia (UFU)

lawag@live.com, mgcarneiro@ufu.br

Introdução: O pôquer, reconhecido internacionalmente como esporte da mente, é amplamente estudado por envolver raciocínio estratégico, tomada de decisão sob incerteza e análise probabilística. Além de sua relevância acadêmica, o jogo oferece um ambiente desafiador para avaliar modelos de linguagem em tarefas dinâmicas e de múltiplas etapas. Com o avanço recente dos Large Language Models (LLMs), surge a oportunidade de investigar sua capacidade de recomendar ações alinhadas a princípios estratégicos e fundamentos da Teoria dos Jogos. No entanto, trabalhos existentes, mesmo aqueles que aplicaram ajustes supervisionados, baseiam-se em gerações anteriores de modelos e raramente exploram abordagens integradas que combinem inferência determinística e prompts estruturados. Essas lacunas motivaram o desenvolvimento de um processo de ajuste supervisionado voltado a aprimorar a consistência, explicabilidade e reprodutibilidade das recomendações. **Objetivos:** Este estudo tem como objetivo avaliar o impacto do fine-tuning supervisionado no desempenho de LLMs abertos aplicados à recomendação de ações estratégicas no pôquer. Busca-se verificar em que medida o ajuste supervisionado aumenta a precisão e estabilidade das decisões geradas, ao mesmo tempo em que aprimora a coerência e clareza das justificativas fornecidas pelos modelos. **Método:** Foram construídas instruções textuais representando cenários reais de jogo, com informações sobre posições, stacks e ações dos participantes. O conjunto de dados foi balanceado e rotulado, reduzindo ambiguidades e garantindo diversidade de contextos. O pipeline desenvolvido integra prompts estruturados, inferência determinística token a token e avaliação por log-probabilidade, compondo o processo de fine-tuning supervisionado. A abordagem é comparada a métodos few-shot, permitindo mensurar ganhos de desempenho e estabilidade. **Resultados:** Em comparação a abordagens publicadas com modelos anteriores ou técnicas heurísticas, o método proposto apresentou ganhos substanciais em precisão e estabilidade. Após o fine-tuning, os modelos alcançaram acurácia próxima de 85%, representando aumento de aproximadamente 13% em relação à configuração few-shot. Além da melhoria quantitativa, observou-se maior consistência entre previsões semelhantes e justificativas mais alinhadas a princípios estratégicos do domínio. Esses resultados evidenciam o potencial do pipeline proposto em reduzir variabilidade e ampliar a confiabilidade das recomendações. **Conclusão:** Os resultados indicam que modelos de linguagem modernos, quando submetidos a ajustes supervisionados e avaliados por métricas determinísticas, podem atuar como assistentes estratégicos em domínios complexos como o pôquer. A metodologia proposta demonstra reprodutibilidade, utiliza exclusivamente modelos open source e abre caminho para a criação de sistemas especializados capazes de apoiar decisões em contextos de alta complexidade e incerteza.

Trilha: Trabalho de Pós-graduação.

Estudo de Destilação de Conhecimento para Classificação de Carcinoma Oral de Células Escamosas em Imagens Histológicas

Carlos A. A. Júnior, Marcelo Z. Nascimento, Thiago P. Ribeiro

Universidade Federal de Uberlândia (UFU)

{carlos.alberto, marcelo.nascimento, tpribeiro}@ufu.br

Introdução: O diagnóstico de Carcinoma Oral de Células Escamosas (OSCC) via imagens histológicas é um desafio importante na classificação de imagens médicas. Diante da necessidade da construção de modelos mais leves em processo de classificação, técnicas como Destilação de Conhecimento (KD) têm ganhado relevância. **Objetivos:** Essa pesquisa busca investigar abordagens de KD para gerar um modelo, originado de uma CNN e baseada em ensemble learning com Stochastic Gradient Descent with Warm Restarts (SGDRE). O estudo explora 2 cenários na abordagem KD, sendo o primeiro experimento a abordagem simples, onde o modelo “aluno”, ao aprende com o melhor “professor” gerado por SGDRE, e um segundo experimento no qual o “aluno” aprende com a média ponderada dos 3 melhores professores, que são ranqueados com a métrica AUC, sendo este o método Ensemble Knowledge Distillation (E-KD). Ambos os métodos também foram comparados com a CNN gerada pelo SGDRE, sem aplicar KD. **Método:** O estudo utilizou um conjunto de dados de 528 imagens histológicas de 100x, composto por 439 imagens de OSCC e 89 imagens normais. Duas arquiteturas de CNN foram definidas, sendo a rede Professor (usada pelo SGDRE), um modelo com quatro camadas convolucionais (32, 64, 128 e 128 filtros) seguido por camadas de max-pooling, uma camada densa de 512 neurônios e uma camada de saída. Esta arquitetura totaliza 6.795.394 parâmetros. A rede Aluno (usada pelo KD) é uma CNN leve, composta por apenas duas camadas convolucionais (32 e 64 filtros) com normalização em lote, seguidas por camada densa de 32 neurônios e camada de saída (4.737.698 parâmetros). A diferença de mais de 2 milhões de parâmetros garante que o Aluno seja um modelo mais leve e rápido para a inferência. Os três métodos comparados, avaliados por validação cruzada de 10 folds, foram: (1) SGDRE (Ensemble), (2) KD (Aluno Simples) e (3) Ensemble-KD. O desempenho foi comparado usando F1-Score. **Resultados:** O método KD (Aluno Simples) demonstrou desempenho superior com F1-Score (0.9285). O método E-KD apresentou F1-Score (0.9149), sendo o segundo melhor. Em contrapartida, o ensemble SGDRE apresentou desempenho inferior aos outros 2 métodos com F1-Score (0.9053). **Conclusão:** O modelo Aluno não apenas foi mais leve, mas também significativamente mais preciso. O modelo E-KD teve desempenho intermediário demonstrando que o modelo ensembles, embora mais complexo, não apresentou resultado superior aos outros experimentos realizados.

Trilha: Trabalho de Pós-graduação.

IA na Educação: Caminhos para uma Transformação Pedagógica Centrada no Ser Humano

Fabienne Charles, Pedro F. Rosa, Shigueo Nomura

Universidade Federal de Uberlândia (UFU)

{fabienne.charles, pfrosi, shigueonomura}@ufu.br

Introdução: Este trabalho analisa como a Inteligência Artificial (IA) pode contribuir para transformar práticas educacionais por meio da personalização do ensino, da ampliação da acessibilidade e da otimização de tarefas pedagógicas. Embora apresente vantagens relevantes, o uso da IA demanda atenção a desafios éticos e estruturais, como privacidade de dados, riscos de viés algorítmico e preservação das interações humanas, essenciais ao aprendizado. Assim, a investigação busca compreender a percepção de diferentes atores educacionais sobre essas tecnologias, bem como identificar oportunidades e limitações para sua integração responsável no ambiente escolar. **Objetivos:** Este estudo tem como objetivo principal avaliar a aceitação, as expectativas e as preocupações relacionadas ao uso da IA na educação, examinando seu potencial para apoiar processos pedagógicos e administrativos. Além disso, busca identificar benefícios percebidos, desafios estruturais e aspectos éticos que devem orientar sua incorporação, justificando a necessidade de políticas e diretrizes que assegurem uma adoção equilibrada e centrada no ser humano. **Método:** A pesquisa adotou uma abordagem qualitativa e exploratória, utilizando um questionário estruturado aplicado a 50 participantes de diferentes perfis educacionais. Foram levantadas percepções sobre conhecimento prévio, experiências com IA, benefícios esperados e desafios associados. A coleta ocorreu de forma on-line, garantindo anonimato e permitindo reunir dados diversificados. A análise seguiu princípios interpretativos, integrando a literatura existente com as opiniões dos respondentes para identificar tendências e pontos críticos. **Resultados:** Os resultados revelam ampla familiaridade com IA e percepções majoritariamente positivas sobre seu potencial educativo. A maioria dos participantes já utilizou ferramentas baseadas em IA e reconhece benefícios como personalização, acessibilidade ampliada e automação de tarefas. Contudo, persistem preocupações com dependência tecnológica, privacidade de dados, falta de interação humana e desigualdade de acesso. A disposição para aprender e adotar IA é elevada, embora acompanhada de cautela quanto a impactos no trabalho docente e à necessidade de capacitação adequada. **Conclusão:** Conclui-se que a IA oferece oportunidades promissoras para aprimorar as práticas educacionais, desde que seja implementada de modo ético, transparente e alinhado às necessidades pedagógicas. O estudo ressalta a importância de preservar o papel humano no ensino e de desenvolver diretrizes que orientem o uso responsável dessas tecnologias. Tais medidas podem favorecer uma integração equilibrada, potencializando benefícios e mitigando limitações, contribuindo para uma educação mais inclusiva, eficaz e adaptada aos desafios contemporâneos.

Trilha: Trabalho de Pós-graduação.

Impacto da Normalização e Redução de Dimensionalidade na Classificação de Imagens Histológicas por Aprendizado de Máquina

Leandro T. T. Oliveira, Marcelo Z. Nascimento

Universidade Federal de Uberlândia (UFU)

leandrotto_1@hotmail.com, marcelo.nascimento@ufu.br

Introdução: Sistemas de apoio ao diagnóstico em patologia digital dependem de pipelines capazes de contribuir com a investigação de lesões em tecidos com a variabilidade de corantes H&E. Variações de coloração, ruído e alta dimensionalidade dos descritores podem degradar a generalização dos classificadores. Embora a literatura reporte benefícios da normalização de corantes e a redução de dimensionalidade, seu efeito combinado sobre descritores texturais clássicos, na classificação binária de lesões orais, permanece pouco quantificado. Este trabalho investigou sistematicamente o impacto dessas etapas no desempenho de modelos supervisionados. **Objetivos:** Avaliar se a normalização de corantes H&E altera a capacidade discriminativa entre tecido saudável e não saudável; Medir o ganho de ensembles por votação frente a classificadores individuais; Identificar combinações de pré-processamento, redução e classificador que maximizem acurácia, recall, precisão, especificidade e AUC. **Método:** Foi elaborado um conjunto de dados obtidos de descritores texturais e não morfológicos extraídos das mesmas ROIs (17 características): uma sem normalização e outra normalizada por estimativa de corantes H&E. O pipeline incluiu análise estatística, detecção/remoção de outliers, padronização e balanceamento com SMOTE. Avaliou-se os modelos SVM, Random Forest, KNN e Naive Bayes sem e com redução de dimensionalidade, variando o número de componentes. Também avaliou a votação majoritária a partir de modelos binários treinados em pares. O desempenho foi medido por acurácia, recall, precisão, especificidade e AUC. **Resultados:** Sem redução de dimensionalidade, os melhores resultados ocorreram na base não normalizada, enquanto a normalização de cor provocou queda consistente em todas as métricas e em todos os classificadores, mais acentuada em Naive Bayes e Random Forest. Com redução de dimensionalidade, observou-se melhora expressiva em ambas as bases; a RPCA foi dominante, elevando acurácia, AUC e especificidade. Na base sem normalização, SVM e RPCA alcançou cerca de 95% de acurácia e AUC próxima de 0,99; na base normalizada, os melhores arranjos atingiram aproximadamente 89% de acurácia e AUC em torno de 0,94. O ensemble por votação apresentou ganhos pequenos a moderados e dependentes da combinação, com maior benefício quando os modelos-base eram complementares. O Naive Bayes foi o classificador mais beneficiado pela redução de dimensionalidade. No entanto, a normalização reduziu a discriminatividade dos descritores manuais, possivelmente por atenuar informações de cromaticidade relevantes. **Conclusão:** A etapa de pré-processamento é determinante para CAD com descritores texturais. A redução de dimensionalidade, especialmente via RPCA, mostrou-se recomendável por aumentar robustez e acurácia. A normalização H&E deve ser aplicada, pois pode reduzir desempenho quando a representação depende de variações cromáticas das cores presentes no tecido.

Trilha: Trabalho de Pós-graduação.

Lógica Informal e Pensamento Crítico no Ensino Introdutório de Programação: Proposição e Validação de Modelo Pedagógico em contextos técnicos e de graduação

Firmiano Reis Silva

Universidade Federal de Uberlândia (UFU)

firmiano@ufu.br

Introdução: O ensino introdutório de programação em cursos técnicos e superiores no Brasil lida com elevados índices de reprovação e desistência. Essa questão se relaciona à dificuldade dos alunos em desenvolver raciocínio lógico e crítico, pois o ensino convencional prioriza lógica formal básica, restringindo a solução de problemas complexos que demandam argumentação contextualizada e tomada de decisão crítica. Evidencia-se lacuna entre a teoria sobre lógica informal e pensamento crítico e sua eficácia empírica no ensino de programação para iniciantes. **Objetivos:** O objetivo é propor e comprovar um modelo pedagógico integrado para o ensino de programação introdutória, com ênfase no aprimoramento do raciocínio lógico-crítico em cursos técnicos e de graduação, organizado em três etapas: (1) Revisão Sistemática da Literatura para avaliar o estado atual e identificar lacunas; (2) desenvolvimento de um modelo pedagógico com base teórica; e (3) comprovação empírica da eficácia em contexto técnico. **Método:** O projeto adota abordagem metodológica em três fases. A primeira realizou RSL (PRISMA 2020) em ACM Digital Library, Scopus, Web of Science e IEEE Xplore (2021-2025). A segunda envolverá proposição de um modelo pedagógico integrado, estruturado em três níveis (formalismo necessário, análise e julgamento crítico, metacognição reflexiva), fundamentado em argumentação lógica. A terceira será validação empírica via estudo-piloto no Instituto Federal do Triângulo Mineiro (IFTM), utilizando Design-Based Research (DBR) com ciclos iterativos. **Resultados:** A RSL analisou 209 registros. Após aplicação rigorosa dos critérios de inclusão e exclusão, 61 estudos foram incluídos, representando taxa de inclusão de aproximadamente 29%. A análise revelou padrão crítico: 39% dos estudos reconhecem limitações pedagógicas da lógica formal pura, argumentando ser abstrata e desconectada de problemas reais. Apenas 7% propõem lógica informal como alternativa viável. De forma preocupante, nenhum dos 61 estudos (0%) integra explicitamente lógica informal E pensamento crítico simultaneamente com rigor empírico em contextos técnicos brasileiros. Limitações metodológicas frequentes incluem amostras pequenas (80%), ausência de grupos controle (65%) e instrumentos inadequados para medir raciocínio crítico (70%). Essa dissociação é reveladora: 39% reconhece que lógica formal isolada falha pedagogicamente, mas ausência total de integração (0%) confirma a lacuna central. **Conclusão:** A lacuna crítica identificada é evidente: enquanto 39% dos estudos reconhecem que lógica formal isolada falha pedagogicamente, nenhum dos 61 estudos (0%) integra explicitamente lógica informal E pensamento crítico simultaneamente com rigor empírico em contextos técnicos brasileiros. Este projeto se diferencia por propor e validar empiricamente essa integração inovadora entre o ensino de lógica para programação ao raciocínio crítico.

Trilha: Trabalho de Pós-graduação.

Modelagem Sequencial e Tabular do Crescimento e Produção em Plantios de Eucalipto

Gabriel Bueno¹, Alvaro A. V. Soares¹, Leonardo I. Rodrigues²,
Tassius M. Araújo², Murillo G. Carneiro¹

¹Universidade Federal de Uberlândia (UFU)

²Celulose Nipo-Brasileira S.A

gabrielbueno@outlook.com, {alvaro.soares, mgcarneiro}@ufu.br
{leonardo.rodrigues, tassius.araujo}@cenibra.com.br

Introdução: A previsão do crescimento florestal é essencial para o manejo sustentável de plantações comerciais, permitindo decisões mais assertivas sobre produtividade e planejamento. Modelos tradicionais, como o de Clutter, baseiam-se em relações lineares e em poucas variáveis, o que limita sua capacidade de representar a complexidade e a dinâmica temporal dos ecossistemas. Nesse contexto, abordagens de aprendizado sequencial, como as arquiteturas Long Short-Term Memory (LSTM), Gated Recurrent Unit (GRU) e Transformer, surgem como alternativas capazes de capturar dependências temporais e padrões não lineares em séries históricas de inventário florestal. **Objetivos:** O estudo avaliou o desempenho de modelos de aprendizado sequencial e tabular na previsão do crescimento e da produção volumétrica de eucaliptos, comparando LSTM, GRU, Transformer, MLP, SVM, XGBoost, TabM e TabPFN com o modelo de Clutter. Também foram testadas variações no número de medições utilizadas como entrada e o efeito da inclusão de variáveis adicionais, como mortalidade, diâmetro, altura média e material genético. **Método:** Utilizaram-se dados de inventário florestal contínuo, coletados entre 2013 e 2019 em 2270 parcelas de *Eucalyptus* spp. no Vale do Rio Doce, Minas Gerais. Cada parcela possuía quatro ou mais medições anuais, com idades de 1,5 a 10 anos. Os modelos receberam sequências de medições anteriores para prever o volume subsequente. O treinamento foi realizado com validação cruzada, otimização de hiperparâmetros via Optuna e divisão de dados em 70% para treino e 30% para teste. As métricas avaliadas foram RMSE, MAE e R^2 . **Resultados:** Considerando o mesmo conjunto de variáveis do modelo de Clutter, os modelos sequenciais e tabulares apresentaram melhor desempenho, com destaque para TabPFN e TabM, que obtiveram menores erros médios (MAE de 9,7 e 11,4 m³/ha) e maiores R^2 (até 0,92). Entre os modelos sequenciais, LSTM e GRU alcançaram R^2 próximos de 0,90 e RMSE em torno de 22 m³/ha, enquanto o Transformer apresentou desempenho semelhante, mas com maior sensibilidade à quantidade de observações. MLP, SVM e XGBoost tiveram resultados competitivos, porém inferiores aos modelos sequenciais. Com a adição de variáveis complementares, TabPFN atingiu R^2 de 0,97 e RMSE de 123 m³/ha, seguido por TabM ($R^2 = 0,96$) e LSTM ($R^2 = 0,94$), demonstrando o ganho obtido com o uso de informações adicionais. **Conclusão:** Os resultados evidenciam que modelos de aprendizado sequencial, especialmente LSTM, GRU e Transformer, são altamente eficazes na modelagem do crescimento florestal, superando o modelo de Clutter e algoritmos clássicos. A inclusão de variáveis adicionais e o uso de diferentes comprimentos de sequência ampliaram a precisão das previsões, indicando que a combinação de abordagens sequenciais e tabulares é uma estratégia promissora para o manejo florestal orientado por dados e para o avanço da sustentabilidade produtiva.

Trilha: Trabalho de Pós-graduação.

Modelo de Propagação da COVID-19 baseado em uma Rede de Autômatos Celulares

Beatriz Viera, Luiz G. A. Martins

Universidade Federal de Uberlândia (UFU)

beatrizvieira515@gmail.com, lgamartins@ufu.br

Introdução: A pandemia de COVID-19 provocou impactos significativos na saúde pública e na dinâmica socioeconômica global. O Brasil registrou mais de 39 milhões de casos e 716 mil óbitos até maio de 2025. A complexidade da disseminação da doença exige modelos capazes de integrar a heterogeneidade espacial, temporal e social dos dados epidemiológicos. Este trabalho propõe um modelo híbrido que combina autômatos celulares (ACs), conectados por um grafo geográfico, com algoritmos de aprendizado de máquina para simular a propagação da COVID-19 em municípios de Minas Gerais. **Objetivos:** Desenvolver um modelo epidemiológico que represente com precisão a evolução espaço-temporal da epidemia, incorporando fatores de mobilidade, clima e variáveis demográficas. **Método:** O estudo segue as etapas do processo KDD: extração e tratamento, seleção de atributos, modelagem preditiva e epidemiológica. Foram utilizados dados públicos do Brasil.IO, Ministério da Saúde, IBGE, Google Mobility e INMET referentes a 2020. As bases foram integradas e normalizadas, com criação de médias móveis, taxas de crescimento e indicadores derivados. A seleção de atributos utilizou RFECV, identificando os preditores mais relevantes entre as variáveis. Na modelagem preditiva, foram avaliados algoritmos baseados em árvores de decisão, ajustados por busca em grade, para estimar o número diário de novos casos, cujas previsões definiram as regras de transição dos ACs. Modelo epidemiológico: Cada município foi representado por uma grade bidimensional de agentes nos estados Suscetível (S), Exposto (E), Infectado (I), Recuperado (R) e Morto (D), seguindo a dinâmica SEIRD. As interações locais consideraram a vizinhança de Moore e a mobilidade dos agentes, enquanto as interações intermunicipais foram modeladas por um grafo que conecta cada cidade aos quatro vizinhos mais próximos. **Resultados:** Entre os modelos testados, o CatBoost apresentou o melhor desempenho na média móvel de sete dias ($RMSE = 57,81$; $MAE = 24,56$; $R^2 = 0,61$), superando XGBoost, Random Forest, Extra Trees e Gradient Boosting. **Discussão:** A integração das previsões do CatBoost às regras SEIRD dos autômatos celulares permitiu reproduzir adequadamente as tendências médias e a variação espacial da pandemia, refletindo a influência da conectividade intermunicipal. Em comparação com estudos que utilizam regras fixas e escalas regionais, o modelo proposto se destaca pela granularidade municipal e pela adaptação automática dos parâmetros epidemiológicos via aprendizado de máquina. **Conclusão:** Como continuidade, pretende-se expandir o treinamento para 2022–2023 visando prever 2024, incluir um modelo de óbitos, incorporar regras sazonais e avaliar abordagens temporais mais avançadas para aprimorar a capacidade preditiva.

Trilha: Trabalho de Pós-graduação.

Novo Modelo de Simulação de Dinâmica de Pedestres para Ambientes Externos baseado em Autômatos Celulares

Eduardo Silva, Luiz Martins

Universidade Federal de Uberlândia (UFU)
{eduardocassiano, lgamartins}@ufu.br

Introdução: um dos maiores desafios da simulação de dinâmica de pedestres (SDP) é a replicação de situações reais. Essa dificuldade é agravada à medida que características adicionais são consideradas, por exemplo, propagação de pânico, desorientação das rotas de fuga, dinâmicas do ambiente, dentre outras. Por meio de regras de transição, os Autômatos Celulares (ACs) são estruturas paralelizáveis que podem ser aplicadas à SDP. Contudo, a maioria dessas pesquisas são voltadas para ambientes fechados, onde boa parte das características essenciais desses cenários já foram exploradas. Esta pesquisa elabora um novo modelo de SDP baseado em ACs para ambientes externos, considerando os desafios da geometria do ambiente, interação entre os pedestres e influência de fenômenos dinâmicos adversos, como desorientação e fogo. **Objetivos:** construir um novo modelo de SDP baseado em ACs que estenda o conceito de piso para reprodução das dinâmicas voltadas aos ambientes externos, tais como desorientação na movimentação, alternância de rotas, obstáculos transponíveis e comunicação entre pedestres. Além disso, pretende-se integrá-lo a um modelo de espalhamento de fogo, de modo a representar a influência dessa dinâmica na movimentação dos pedestres. **Método:** o novo modelo baseia-se no piso determinista do modelo Varas, utilizando o cálculo das distâncias e as regras para conflitos de pedestres. Em seguida, nosso modelo também adotou o piso dinâmico do modelo Alizadeh, onde os pedestres podem alternar sua rota de fuga considerando as saídas congestionadas. Logo após, características como desorientação, obstáculos transponíveis e placas de orientação foram acrescentadas. Atualmente, o modelo integra uma dinâmica de incêndio voltada a avaliar os efeitos da comunicação entre pedestres para a sinalização de regiões inseguras. **Resultados:** os dados coletados mostram que a inclusão das novas dinâmicas melhorou significativamente o realismo e o desempenho das simulações. Os obstáculos transponíveis permitem que pedestres encontrem rotas alternativas, acelerando o tempo total de desocupação dos cenários. A inclusão da desorientação faz os pedestres terem dificuldade em encontrar as rotas certas em regiões isoladas do mapa, tornando a simulação mais realista. Em contrapartida, as placas de sinalização reduzem os efeitos da desorientação, agilizando o processo de desocupação. **Conclusão:** o novo modelo estende o conceito de piso para ambientes externos, onde novos desafios são considerados. Nesses cenários, foram acrescentadas dinâmicas para exploração de obstáculos transponíveis, desorientação das rotas de fuga e auxílio por meio de placas de sinalização, aprimorando o realismo e desempenho da SDP. Futuramente, pretende-se melhorar a dinâmica de incêndio com a integração de um modelo também baseado em ACs, além também da inclusão de um mecanismo de comunicação entre os pedestres, visando aperfeiçoar a verossimilhança perante a complexidade desses cenários.

Trilha: Trabalho de Pós-graduação.

Perfil dos Participantes em Treinamentos de Programação Competitiva em uma Instituição de Ensino Superior Moçambicano (2024–2025)

Danilo Richards, Luiz Cláudio Theodoro,
Rafael D. Araújo, João Henrique S. Pereira

Universidade Federal de Uberlândia (UFU)
{danilo.richards, luiz.theodoro, rafael.araujo, joaohs}@ufu.br

Introdução: Na área da Computação, a programação é uma competência fundamental, valorizada tanto na formação acadêmica quanto no mercado profissional. As metodologias ativas baseadas em programação competitiva têm se destacado por integrar teoria e prática na resolução de problemas por meio do código, estimulando o raciocínio lógico, o pensamento computacional e a autonomia dos estudantes. Essas abordagens aproximam o aprendizado de situações reais e favorecem o trabalho colaborativo, a criatividade e a persistência diante de desafios computacionais. Apesar do seu potencial pedagógico, ainda é necessário compreender melhor os efeitos da programação competitiva quando aplicada em contextos educacionais não voltados à competição, especialmente onde há limitações de recursos e de suporte docente. Em Moçambique, essa prática ainda é recente e enfrenta desafios estruturais e metodológicos, reforçando a importância de estudos que contribuam para sua consolidação no ensino superior. **Objetivos:** O estudo teve como objetivo principal compreender quem são os estudantes que participam desses treinamentos, bem como suas motivações, dificuldades e níveis de habilidade técnica, informações fundamentais para o aprimoramento e a sustentabilidade dessas iniciativas. **Método:** A pesquisa adotou uma abordagem descritiva, com base em dados coletados por meio de questionários estruturados aplicados aos participantes de treinamentos realizados em 2024 e 2025. As variáveis analisadas incluíram faixa etária, gênero, ano do curso, motivações, dificuldades enfrentadas, nível de habilidade e frequência de participação. **Resultados:** Os resultados mostram que o público participante é composto, em sua maioria, por estudantes com idades entre 17 e 23 anos, matriculados principalmente no 2º e 3º anos dos cursos de Computação. O nível de habilidade declarado varia de básico a intermediário. As principais motivações relatadas foram o aprimoramento da lógica de programação, o interesse em aprender novas técnicas de algoritmos e o desejo de melhorar o desempenho acadêmico. Entre as dificuldades mais recorrentes, destacam-se a gestão do tempo associada à sobrecarga de responsabilidades acadêmicas com os treinamentos e a ausência de feedback detalhado nas atividades, consequência do uso de juízes online que, ao informarem apenas se a submissão está correta ou incorreta e não explicam a natureza dos erros. **Conclusão:** Os treinamentos de programação mostraram-se uma ferramenta importante para o fortalecimento das competências lógicas e analíticas para os participantes. Os desafios relacionados à gestão do tempo, permanecem como o principal entrave à participação contínua. Ainda existem etapas e desafios a serem transpostos para o aperfeiçoamento da metodologia, adaptada às condições locais, e para a consolidação da programação competitiva como prática formativa e para a implementação de competições no contexto moçambicano.

Trilha: Trabalho de Pós-graduação.

Predição de Desempenho Acadêmico Discente Utilizando Dados de Frequência em Aulas: Uma Abordagem com Múltiplos Modelos de Aprendizado de Máquina

Gustavo Henrique R. Magalhães,
Paulo Henrique R. Gabriel, José Gustavo S. Paiva

Universidade Federal de Uberlândia (UFU)
{gustavohrm, phrg, gustavo}@ufu.br

Introdução: A predição de desempenho acadêmico é um desafio central na Mineração de Dados Educacionais. A literatura foca em dados demográficos, notas prévias ou interações em Sistemas de Gerenciamento de Aprendizagem, como cliques ou logins, tratando a frequência de forma puramente quantitativa. Este trabalho aborda uma lacuna: o valor preditivo do conteúdo textual específico das aulas em que o aluno esteve ausente. A hipótese é que o quê foi perdido é mais informativo do que quanto foi perdido. **Objetivos:** O objetivo é avaliar a viabilidade de prever o desempenho (conceitos 'A' a 'E') usando Processamento de Linguagem Natural aplicado ao conteúdo das aulas ausentes. Buscamos também identificar os melhores algoritmos e vetorizadores para esta nova representação textual. A justificativa é a necessidade de criar sistemas de alerta precoce mais eficazes, que ofereçam um diagnóstico (o porquê do risco, ou seja, os tópicos perdidos), superando métodos que apenas indicam que o aluno está em risco. **Método:** Diferente de trabalhos que focam em dados de Sistemas de Gerenciamento de Aprendizagem, o método construiu um atributo preditivo a partir do registro de conteúdo de aulas de um conjunto de dados da FACOM/UFU (2011-2019). Para cada aluno, um atributo de texto foi gerado agregando o conteúdo das aulas em que esteve ausente. O texto foi processado por CountVectorizer e TfidfVectorizer. Oito algoritmos de classificação (incluindo Random Forest, CatBoost e XGBoost) foram treinados e avaliados. A performance foi medida por métricas ponderadas (F1-Score, Recall) devido ao desbalanceamento natural das classes de conceito (A-E). **Resultados:** O desempenho agregado (F1-Score Ponderado) foi modesto (~ 0.44), com a escolha do algoritmo (Random Forest, CatBoost) sendo mais determinante que a do vetorizador. A análise granular por classe revelou o potencial do método: a predição da Classe 'E' (reprovação) foi notavelmente alta (Recall médio de 0.54 geral, 0.64 em disciplinas GSI, e picos de $F1 > 0.90$ em algumas disciplinas). Em contraste, a Classe 'D' (risco intermediário) mostrou-se muito difícil de prever. A discussão sugere que o modelo é mais eficaz em disciplinas especializadas, onde o conteúdo é mais estável (levantando a hipótese de menor rotatividade docente), em oposição a disciplinas de ciclo básico com maior ruído textual (maior rotatividade). **Conclusão:** Conclui-se que a abordagem não é um preditor universal de notas (dada a dificuldade nas classes A, B, D), mas valida-se como uma ferramenta robusta e viável para sistemas de alerta precoce. A capacidade de identificar com alta precisão e recall os alunos em trajetória de reprovação (Classe 'E') oferece valor prático. A contribuição principal é fornecer um diagnóstico inicial (os tópicos perdidos), uma vantagem sobre métodos tradicionais que usam apenas a contagem de frequência.

Trilha: Trabalho de Pós-graduação.

Projeto e Verificação da Olimpíada Brasileira de Programação Quântica

Giullia Menezes, João Henrique Pereira, Rafael Araújo

Universidade Federal de Uberlândia (UFU)

{giullia.rodrigues, joaohs, rafael.araujo}@ufu.br

Introdução: A área da computação quântica tem apresentado evoluções significativas a partir do desenvolvimento das arquiteturas de hardware nos últimos anos, isto tem possibilitado implementar em computadores quânticos reais algoritmos que eram apenas teóricos. **Objetivos:** Dado o potencial da área de computação quântica e a relevância das competições científicas e maratonas de programação para o sistema educacional de computação para estimular discentes à maior dedicação e desempenho nos estudos, o objetivo geral desta pesquisa é implementar e validar a Olimpíada Brasileira de Programação Quântica (OBPQ), dentro das diretrizes da Diretoria de Competições Científicas da Sociedade Brasileira de Computação (SBC). **Método:** Para isto, será necessário compreender as diretrizes da SBC para competições científicas e, também, as características das arquiteturas atuais de computação quântica para definir aquela que se adequa ao contexto brasileiro a ser utilizada na OBPQ. **Resultados:** Até o momento, foi feito o mapeamento do estado da arte das tecnologias de computação quântica e suas características tecnológicas. O Qiskit é o conjunto de softwares mais popular e eficiente do mundo para computação quântica e pesquisa de algoritmos. A plataforma QCoder se mostrou uma ótima plataforma tecnológica para administração automatizada da competição, nela é possível acessar exercícios que exploram diversos conceitos importantes da computação quântica. A computação quântica ainda está sendo explorada e não está presente em cursos de graduação, diante disso, entende-se que é necessária uma explicação prévia de conceitos importantes que serão trabalhados ao longo da competição. Neste momento estão sendo levantados materiais que explorem esses conceitos e confecção de exercícios para a olimpíada. Simultaneamente está sendo feito o contato com membros do Grupo de Interesse de Computação e Comunicação Quântica (GI-CCQ) e Diretoria de Competições Científicas da Sociedade Brasileira de Computação (SBC) para levantar informações que são importantes na competição, considerando a visão de especialistas. **Conclusão:** Uma das maiores contribuições científicas deste trabalho é, incentivar discentes a estudar computação quântica de alto desempenho por meio da Olimpíada Brasileira de Programação Quântica.

Trilha: Trabalho de Pós-graduação.

Técnicas de IA Explicável aplicadas ao Problema de Manutenção Preditiva na Indústria Alimentícia

Higor Freitas, Flavio S. Silva, Fabíola Pereira

Universidade Federal de Uberlândia (UFU)

{higor.freitas, flavio, fabiola.pereira}@ufu.br

Introdução: Este trabalho propõe a aplicação integrada de Manutenção Preditiva (PdM) e Inteligência Artificial Explicável (XAI) para reduzir significativamente o descarte de produtos no processo de cozimento de mortadela. As perdas, denominadas “quebras”, classificam-se em dois tipos: quebra por baixo, resultante de cozimento insuficiente, e quebra por alto, decorrente de cozimento excessivo, situação em que o produto se torna impróprio para consumo. A motivação do projeto surge do índice insatisfatório de perdas observado na empresa, causado por tempo de cozimento inadequado, distribuição desigual de calor nas estufas e dificuldade de análise em tempo real devido à coleta manual de dados. **Objetivos:** O objetivo geral é desenvolver um modelo de IA explicável para identificar e reduzir quebras na produção de mortadela defumada, validando o processo decisório dos algoritmos utilizados. Os objetivos específicos incluem o pré-processamento dos dados de produção, o treinamento de modelos de machine learning, a avaliação de desempenho, a aplicação de técnicas de XAI pós-hoc e a comparação de abordagens de explicabilidade para suporte à tomada de decisão industrial. **Método:** A solução desenvolvida combina modelos preditivos, como Redes Neurais Profundas (DNNs) e Random Forest, com técnicas de explicabilidade como LIME e SHAP, visando não apenas prever falhas, mas também interpretar as decisões do modelo, conferindo transparência e confiança ao processo. A fundamentação teórica baseia-se na Manutenção Preditiva, que utiliza dados em tempo real e algoritmos de machine learning para prever falhas antes de sua ocorrência, e na Inteligência Artificial Explicável, que busca aprimorar a transparência e a interpretabilidade dos modelos de IA. O estudo envolveu coleta e normalização de dados a partir de 2.622 registros de quebra e 4.814 registros de estufa, resultando em 2.920 linhas consolidadas após o pré-processamento. **Resultados:** Os modelos de DNN e Random Forest alcançaram acurácias de 95% e 90%, respectivamente, com foco principal na aplicação de técnicas pós-hoc para interpretação dos resultados. **Conclusão:** Os resultados esperados incluem previsão precisa de quebras com base em temperatura, tempo e posicionamento, explicações claras sobre como as variáveis influenciam as falhas, redução de perdas e custos operacionais, além do desenvolvimento de um framework replicável para outros processos industriais.

Trilha: Trabalho de Pós-graduação.

Um panorama da Infraestrutura Tecnológica para Implementação da Programação Competitiva em Moçambique

Danilo Richards, Luiz Cláudio Theodoro,
Rafael D. Araújo, João Henrique S. Pereira

Universidade Federal de Uberlândia (UFU)
{danilo.richards, luiz.theodoro, rafael.araujo, joaohs}@ufu.br

Introdução: A infraestrutura tecnológica é um componente essencial nas instituições de ensino superior, sobretudo nos cursos ligados à computação. Laboratórios de informática, conectividade estável, servidores funcionais e políticas de manutenção contínua são elementos centrais para sustentar o ensino, a pesquisa e a extensão voltadas à inovação. A adoção de metodologias ativas, como a programação competitiva, requer condições técnicas mínimas que garantam o funcionamento dos ambientes de desenvolvimento e avaliação automática. Nesse contexto, compreender a realidade das instituições moçambicanas torna-se fundamental para identificar as potencialidades e limitações que influenciam a viabilidade de implementação dessa abordagem no ensino superior. **Objetivos:** Este estudo tem como objetivo analisar a infraestrutura tecnológica das instituições de ensino superior moçambicanas, considerando aspectos físicos, técnicos e organizacionais, com vistas a compreender o nível de preparação institucional e identificar os principais desafios estruturais que podem afetar a viabilidade de sua implementação. **Método:** Adotou-se uma abordagem descritiva e diagnóstica, utilizando questionários estruturados e entrevistas semiestruturadas com diretores de curso e gestores de tecnologia da informação de instituições de ensino superior moçambicanas. Os dados foram analisados de forma quantitativa e qualitativa, permitindo comparar as instituições quanto à disponibilidade e ao desempenho de seus recursos tecnológicos. **Resultados:** Os dados revelam disparidades na infraestrutura tecnológica entre as instituições analisadas. Verificou-se variação no número de laboratórios, na quantidade de computadores operacionais e na estabilidade da conexão à internet. Em muitas instituições, o suporte técnico é limitado e as políticas de manutenção preventiva são inexistentes ou pouco sistemáticas. Embora a maioria possua laboratórios de informática, a conectividade instável, a manutenção irregular e a escassez de equipamentos operacionais permanecem como desafios recorrentes que restringem o pleno aproveitamento dos recursos disponíveis. **Conclusão:** Constata-se que as instituições de ensino superior moçambicanas enfrentam limitações estruturais importantes que podem dificultar a adoção da programação competitiva em larga escala. A relação entre o número de computadores e o total de estudantes é geralmente desfavorável, e muitos laboratórios são compartilhados entre diferentes cursos ou utilizados como salas de aula, reduzindo ainda mais sua disponibilidade para práticas de programação. Embora haja maior consciência sobre a importância da infraestrutura digital, ela raramente se converte em investimentos concretos. O estudo aponta a necessidade de ações sustentáveis para ampliar e qualificar a infraestrutura, fortalecendo a formação técnica e científica no ensino superior moçambicano.

Trilha: Trabalho de Pós-graduação.

Uma Nova Arquitetura Híbrida para Otimização Multiobjetivo de Hiperparâmetros em IDS

Salomão Pena, Silvio E. Quincozes, Shigueo Nomura

Universidade Federal de Uberlândia (UFU)

{salomao.pena, sequincozes, shigueonomura}@ufu.br

Introdução: o cenário tecnológico atual, marcado pela crescente adoção de dispositivos conectados à Internet das Coisas (IoT), tem resultado na expansão de ataques cibernéticos, exigindo soluções de segurança inteligentes e adaptáveis para a detecção de intrusões. A complexidade e a sofisticação das ameaças cibernéticas impulsionam a busca por Sistemas de Detecção de Intrusões (IDS) capazes de monitorar e responder aos incidentes cibernéticos em tempo real. **Objetivos:** o principal objetivo desta pesquisa é propor e validar uma nova arquitetura para a otimização multiobjetivo de hiperparâmetros em IDS baseados em aprendizado por reforço profundo. Para atingir o objetivo, desenvolveu-se uma nova arquitetura híbrida, denominada SAC&LSTM (Soft Actor-Critic integrado com redes Long Short-Term Memory), que utiliza Otimização Bayesiana Multiobjetivo (BO), por meio da biblioteca Optuna, para realizar o ajuste automático de hiperparâmetros críticos, incluindo o número de camadas, a taxa de aprendizado, a entropia e a função de ativação. **Método:** o trabalho foi conduzido como uma pesquisa experimental e quantitativa, avaliando o desempenho da arquitetura SAC&LSTM+BO em um ambiente simulado de detecção de intrusão. A escolha da arquitetura híbrida foi motivada pela necessidade de combinar a capacidade das redes LSTM de modelar sequências temporais e dependências de longo prazo em dados de tráfego de rede com a eficiência do algoritmo SAC, robusto para espaços de ação contínuos e otimiza a máxima entropia, promovendo uma exploração mais eficaz de ações. **Resultados:** os resultados experimentais evidenciaram o alto desempenho da arquitetura proposta na detecção de intrusões. O agente SAC&LSTM otimizado alcançou um desempenho notável no conjunto de teste, com F1-score superior a 99,42% e acurácia de 98,89%. A precisão de 99,69% e o recall de 99,16% atestam a robustez do modelo, mesmo em cenários de classes desbalanceadas. Em relação aos fatores de custo computacional, a otimização multiobjetivo provou-se extremamente eficiente. Os resultados médios indicaram um consumo máximo de memória RAM de 203,96 MB e um tempo médio de otimização por época de 11,08 minutos. Estes valores mostram que a arquitetura é viável para aplicações em tempo real e em ambientes com restrição de recursos, como a IoT, uma limitação frequentemente ignorada em trabalhos correlatos. **Conclusão:** este estudo propôs e validou a arquitetura SAC&LSTM integrada com BO para aprimorar a detecção de intrusão em redes de IoT. A utilização do conjunto de dados realista CICIoT2023 garantiu que os resultados fossem representativos dos desafios de segurança modernos. As principais contribuições deste trabalho residem na proposta de uma arquitetura híbrida de aprendizado por reforço profundo para IDS, na validação em um cenário IoT contemporâneo com múltiplas métricas relevantes, incluindo fatores de custo, e na definição de uma função de recompensa para otimização multiobjetivo.

Trilha: Trabalho de Pós-graduação.

Visualizações em Vídeos de Vigilância Orientadas por Ontologia

João Pedro Nunes, José Gustavo Paiva

Universidade Federal de Uberlândia (UFU)

{joao.nunes2, gustavo}@ufu.br

Introdução: Sistemas de videomonitoramento geram grandes volumes de dados, tornando inviável a análise exclusivamente humana de longas gravações, o que dificulta a identificação efetiva de eventos relevantes. Muitas abordagens computacionais atuais executam diversas tarefas analíticas, como detecção, sumarização e predição de eventos, que muitas vezes não oferecem uma relação semântica ao contexto de segurança, criando uma distância entre esse processo e as reais necessidades dos analistas. As ontologias são artefatos de representação do conhecimento, estruturadas em domínios específicos que podem ser compreendidas por máquinas dada sua formalização em linguagem OWL (Ontology Web Language). Estas possuem relações semânticas estruturadas em triplas Sujeito - Predicado (ação) - Objeto, com as quais são criados axiomas para representar domínios específicos do conhecimento com regras e com riqueza semântica. **Objetivos:** Criar uma abordagem de análise visual de vigilância baseada na descrição do domínio do conhecimento fornecido por uma ontologia, de modo a alcançar representações que tornem gravações extensas mais inteligíveis para análise, priorizando relações e contextos relevantes. **Método:** O projeto será desenvolvido em três etapas principais: (1) construção de um módulo de detecção automática de relações entre objetos identificados em vídeo de vigilância; (2) construção de um módulo de instanciação dos eventos na ontologia, preservando informações espaço-temporais; (3) implementação de um módulo de análise visual da ontologia produzida com foco na exploração semântica dos eventos e relações entre objetos. **Resultados:** A avaliação envolverá estudos de caso em vídeos de vigilância diversos, analisando a capacidade da proposta no conteúdo desses vídeos. Espera-se obter uma ferramenta eficaz de construção de ontologias e sua compreensão em termos de conteúdo, uma ontologia aplicada ao domínio de vigilância, um protótipo funcional de sumarização e navegação semântica de vídeos e evidências de que a combinação entre extração automática baseada em biblioteca especializada e representação ontológica favorece a identificação rápida e explicável de situações de interesse. **Conclusão:** A proposta encontra-se em estágio inicial de desenvolvimento e estão sendo feitos testes em bibliotecas de detecção de eventos em vídeos que serão utilizadas nas visualizações futuras.

Trilha: Trabalho de Pós-graduação.