

36º CONGRESSO DE INICIAÇÃO CIENTÍFICA E TECNOLÓGICA DA UFRN – CICT 2025



36º cict

Planeta Água: a cultura oceânica
para enfrentar as mudanças
climáticas no meu território.

ANAIS

EXPEDIENTE

APRESENTAÇÃO

O Congresso de Iniciação Científica e Tecnológica (CICT) é um evento aberto à comunidade no qual todos os aproximadamente 1.700 alunos de Iniciação Científica e Tecnológica da UFRN (bolsistas e voluntários) apresentam os resultados de suas pesquisas desenvolvidas ao longo de um ano, como cumprimento de um plano de trabalho elaborado e orientado por um professor/pesquisador do quadro permanente da Instituição.

O CICT está entre as ações inseridas no âmbito do Programa Institucional de Bolsas de Iniciação Científica e Tecnológica da UFRN, que possui entre seus objetivos: a) Contribuir para a formação e engajamento de recursos humanos em atividades de pesquisa, desenvolvimento tecnológico e inovação; b) Contribuir para a formação científica de recursos humanos que se dedicarão a qualquer atividade profissional; e c) Contribuir para a formação de recursos humanos que se dedicarão ao fortalecimento da capacidade inovadora das empresas no País.

A 36ª edição do Congresso teve como tema “Planeta Água: a cultura oceânica para enfrentar as mudanças climáticas no meu território”, estando alinhada aos 17 Objetivos do Desenvolvimento Sustentável da Agenda 2030 da ONU. Ao submeter os trabalhos ao Congresso, os alunos indicaram os objetivos da Agenda 2030 aos quais suas pesquisas estavam relacionadas. Essa iniciativa visa estimular a pesquisa científica e tecnológica voltada para a busca por soluções para os desafios sociais contemporâneos. Em uma etapa inicial do CICT, denominada Mostra de Trabalhos e Vídeos de Ciência, Tecnologia e Inovação, foram realizadas as apresentações dos trabalhos submetidos pelos participantes do evento. As apresentações ocorreram de forma assíncrona, por meio de arquivos digitais (trabalho completo e vídeo), sendo disponibilizadas não só aos avaliadores do congresso pelos Sistemas da UFRN, mas também ao público em geral por meio do site do evento (www.cic.propesq.ufrn.br). Após essa etapa inicial, que ocorreu sob o formato virtual, os 48 melhores trabalhos do evento foram selecionados para apresentação oral em uma etapa complementar do congresso que ocorreu de forma presencial.

Os alunos concorreram ao 9º Prêmio Destaque na Iniciação Científica e Tecnológica da UFRN, que foi instituído pela Pró-Reitoria de Pesquisa em 2017 para reforçar as ações dos programas institucionais de iniciação científica e tecnológica. O prêmio contemplou as categorias Trabalho Destaque de Iniciação Científica e Tecnológica e Vídeo Destaque de Iniciação Científica e Tecnológica) e na subcategoria de Inovação. Os premiados foram agraciados com bolsas mensais de até R\$ 1100,00 (mil e cem reais) durante o período de um ano. Os orientadores dos alunos premiados adquiriram precedência em relação aos demais para efeitos de concorrência no Edital de Bolsas de Pesquisa da UFRN em 2026.

Ainda em 2025, a Pró-Reitoria de Pesquisa realizou a 7ª Edição do Prêmio Pesquisador Destaque da UFRN, certame instituído em 2019 e que é concedido anualmente aos pesquisadores da instituição que tenham apresentado relevantes contribuições para o desenvolvimento da ciência em cada uma das três grandes áreas do conhecimento: Ciências da Vida; Ciências Exatas, da Terra e Engenharias; e Ciências Humanas e Sociais, Letras e Artes. Além de troféu e certificado, todos os premiados foram contemplados com auxílio no valor de R\$ 10.000,00 (dez mil reais) para custear despesas específicas, discriminadas em edital, e tiveram a oportunidade de compartilhar suas experiências com os participantes do CICT 2025 em uma mesa redonda.

O congresso cumpriu três funções valiosas para o desenvolvimento da ciência na instituição: inserção dos discentes em um ambiente acadêmico de publicação e apresentação dos resultados da pesquisa; divulgação científica por meio dos vídeos que utilizam uma linguagem científica, porém acessível a todos; e, por fim, a popularização da ciência.

PROGRAMAÇÃO

Outubro – 2025

15/10 Divulgação da lista dos pré-selecionados para concorrer ao 9º Prêmio Destaque na Iniciação Científica e Tecnológica da UFRN

Novembro – 2025

03/11 Abertura oficial do evento

03/11 a 07/11 CICT 2025 – Etapa presencial

04/11 a 06/11 Sessões de Apresentações Orais

07/11 Cerimônia de premiação e encerramento do evento

IMD

INSTITUTO METRÓPOLE DIGITAL

CÓDIGO: ET0018

AUTOR: GABRIEL LIMA VIDAL

ORIENTADOR: HEITOR MEDEIROS FLORENCIO

TÍTULO: Implantação de uma bancada de teste para controladores industriais.

Resumo

Este trabalho apresenta a implantação de uma bancada de testes utilizando o CLP XP340, com foco na integração entre controladores industriais e tecnologias de Internet das Coisas (IoT). O projeto contempla o monitoramento do nível de um tanque e o controle remoto via protocolo MQTT e plataforma Node-RED. A proposta visa criar um ambiente didático e funcional para simular processos industriais, promovendo o aprendizado prático na área de automação. Os resultados demonstram que a comunicação entre os dispositivos é eficiente e adequada para aplicações industriais e educacionais.

Palavras-chave: CLP; XP340; IoT; MQTT; Node-RED; Automação industrial;

TITLE: Implementation of a test bench for industrial controllers.

Abstract

This work presents the implementation of a test bench using the XP340 PLC, focusing on the integration of industrial controllers and Internet of Things (IoT) technologies. The project includes tank level monitoring and remote control via the MQTT protocol and the Node-RED platform. The proposal aims to create a functional and educational environment to simulate industrial processes, promoting practical learning in the automation field. The results demonstrate that communication between the devices is efficient and suitable for industrial and educational applications.

Keywords: PLC; XP340; IoT; MQTT; Node-RED; Industrial Automation.

Introdução

O avanço da Internet das Coisas (IoT) tem impactado diretamente os sistemas de automação industrial, exigindo a integração entre controladores programáveis e plataformas digitais. Este trabalho apresenta o desenvolvimento de uma bancada de testes com o CLP XP340, integrando sensores e comunicação remota via protocolo MQTT com a plataforma Node-RED. A proposta visa simular, de forma prática, um sistema de supervisão e controle em tempo real, destacando a viabilidade de sua aplicação em ambientes industriais. O foco da pesquisa está na construção do painel elétrico, cabeamento dos dispositivos e configuração do fluxo de mensagens entre os elementos do sistema.

Metodologia

Foi utilizado o CLP XP340 para o controle do nível de um tanque de água. O sistema foi conectado a sensores analógicos e uma bomba analógica, sendo configurado para enviar os dados coletados ao Node-RED instalado em um computador. A comunicação foi feita via protocolo MQTT, utilizando tópicos para envio dos valores lidos e recepção do setpoint. O Node-RED, por sua vez, processa os dados e permite o ajuste do nível desejado e a visualização de gráfico remotamente. A bancada foi montada no ConecteLab, localizado no prédio do NPITI (IMD) da UFRN, com cabeamento, painel elétrico e estrutura de testes desenvolvidos pelo autor.

Resultados e Discussões

Durante os testes, observou-se que a leitura do nível do tanque foi transmitida em tempo real ao painel no Node-RED. A cada alteração no valor do setpoint, o CLP recebia a mensagem MQTT e ajustava o controle de saída do sistema, promovendo uma resposta imediata. A integração mostrou-se estável e responsiva. Essa abordagem proporciona uma visão aplicada sobre tecnologias emergentes como IoT na automação industrial.

Conclusão

O trabalho alcançou com êxito a integração entre o CLP XP340 e a plataforma Node-RED por meio do protocolo MQTT, permitindo o envio e recebimento de dados em tempo real. A bancada demonstrou desempenho estável e funcional, reproduzindo com fidelidade um cenário de automação industrial com supervisão remota. Como destaque, o sistema permitiu o monitoramento contínuo do nível de um tanque e o ajuste remoto de setpoints, validando a eficiência da comunicação bidirecional.

Referências

- (1) TI, L. Altus Sistemas de Automação | Altus. Disponível em: <<https://www.altus.com.br/base-conhecimento/categoria/19/detalhe/494/saiba-como-utilizar-o-protocolo-mqtt-na-serie-nexto>>. Acesso em: 22 abr. 2025.
- (2) TI, L. Base de conhecimento | Altus. Disponível em: <<https://www.altus.com.br/base-conhecimento/busca?palavrachave=&validadorantispam=6500931>>. Acesso em: 22 abr. 2025.

Anexos

Bancada de teste

MPS-PA

Comunicação MQTT

CÓDIGO: ET0399

AUTOR: SAULO CARLOS PEREIRA DA ROCHA

COAUTOR: ANDERSON CLAUDIO RODRIGUES DA SILVA

COAUTOR: RODRIGO FERREIRA DA NOBREGA

ORIENTADOR: ROGER KREUTZ IMMICH

TÍTULO: Autenticação Descentralizada de Dispositivos IoT com Smart Contracts em Blockchain

Resumo

A crescente interconexão de dispositivos na Internet das Coisas (IoT) impõe desafios de autenticação e segurança que não são plenamente atendidos por arquiteturas centralizadas. Este trabalho apresenta o desenvolvimento de uma arquitetura de contrato inteligente minimamente viável (MVSC) para registro e autenticação descentralizada de dispositivos IoT em blockchain. Foram implementados quatro contratos independentes, voltados ao gerenciamento de sensores de temperatura, umidade, proximidade e movimento, incorporando mecanismos de controle de acesso, prevenção de duplicidade e verificação temporal de validade. A prototipagem ocorreu em Solidity, com validação no ecossistema Ethereum, utilizando Remix, Hardhat e Ganache. Os testes automatizados confirmaram a eficácia dos contratos em cenários de registros válidos, bloqueio de duplicidades, rejeição de parâmetros inválidos, autenticação confiável e restrição de registros por contas não autorizadas. Os resultados demonstram a viabilidade da integração entre blockchain e IoT, estabelecendo uma base para soluções mais robustas, com perspectivas de evolução para funcionalidades avançadas de segurança.

Palavras-chave: IoT, Blockchain, Contratos Inteligentes, Autenticação de Dispositivos

TITLE: Decentralized Authentication of IoT Devices with Smart Contracts on Blockchain

Abstract

The growing interconnection of devices in the Internet of Things (IoT) poses authentication and security challenges that are not fully addressed by centralized architectures. This paper presents the development of a minimally viable smart contract (MVSC) architecture for decentralized registration and authentication of IoT devices on a blockchain. Four independent contracts were implemented, aimed at managing temperature, humidity, proximity, and motion sensors, incorporating access control mechanisms, duplication prevention, and temporal validity verification. Prototyping was carried out in Solidity, with

validation in the Ethereum ecosystem using Remix, Hardhat, and Ganache. Automated testing confirmed the effectiveness of the contracts in scenarios involving valid registrations, blocking duplicates, rejecting invalid parameters, reliable authentication, and restricting registrations by unauthorized accounts. The results demonstrate the feasibility of integrating blockchain and IoT, laying a foundation for more robust solutions, with prospects for evolving into advanced security features.

Keywords: IoT, Blockchain, Smart Contracts, Device Authentication

Introdução

A Internet das Coisas (IoT) tem se consolidado como um dos principais vetores de transformação digital, promovendo a integração de dispositivos heterogêneos em domínios como saúde, transporte, indústria e cidades inteligentes [Majd and Sharko 2023, Garba et al. 2023]. Entretanto, a crescente interconexão desses dispositivos amplia a superfície de ataque e expõe vulnerabilidades relacionadas à privacidade, integridade dos dados e confiabilidade das comunicações [Hameed et al. 2021, Qiu et al. 2024]. Esses desafios evidenciam a necessidade de mecanismos descentralizados, transparentes e resistentes a adulterações, capazes de sustentar a autenticação e o controle de acesso em ambientes distribuídos [Hameed et al. 2021].

Dentre os sistemas de segurança demandados pela IoT, a autenticação de dispositivos se destaca como requisito fundamental, pois assegura a verificação da identidade antes da concessão de acesso e constitui a primeira linha de defesa contra operações não autorizadas [Liou and Lin 2021]. Todavia, soluções tradicionais ainda se apoiam em arquiteturas centralizadas, sujeitas a vulnerabilidades decorrentes do ponto único de falha, além de restrições no compartilhamento seguro de informações [Hameed et al. 2021, Khashan and Khafajah 2023, Qiu et al. 2024]. Nesse contexto, a tecnologia blockchain surge como alternativa promissora, ao combinar criptografia, consenso distribuído e redes peer-to-peer para prover descentralização, imutabilidade, transparência e rastreabilidade [Goswami and Choudhury 2022, Liu et al. 2024]. Essas propriedades contribuem para superar as fragilidades dos mecanismos convencionais e reforçar a confiabilidade dos processos de autenticação em ambientes heterogêneos [Akkaoui 2021, Bargayary and Medhi 2023].

Nesse cenário, os contratos inteligentes (*smart contracts*) assumem papel central ao fornecer os mecanismos necessários para a descentralização e automação dos processos de registro e autenticação dos dispositivos. Em termos gerais, os smart contracts são trechos de código responsáveis por definir as partes envolvidas e as condições necessárias para a execução de transações e interações automatizadas. Devido à sua flexibilidade funcional, trabalhos como os apresentados por [Hussein et al. 2024], [Mahmoud et al. 2024] e [Majd and Sharko 2023] demonstram a aplicabilidade desse recurso em múltiplas frentes, abrangendo o armazenamento de *logs* de autenticação, o gerenciamento da emissão, validação e revogação de *tokens*, além da verificação da autenticidade dos dispositivos antes de sua integração à rede. Nessas propostas, os contratos inteligentes configuram-se como elemento-chave para garantir a confiabilidade e a segurança no provisionamento dos equipamentos conectados.

Diante desse contexto, este trabalho busca desenvolver uma arquitetura descentralizada para o gerenciamento e controle de *smart contracts* voltados ao registro e à autenticação de dispositivos IoT em rede blockchain. A proposta consiste em implementar um protótipo que

permita o cadastramento de dispositivos por um administrador autorizado e a verificação de sua autenticidade por qualquer entidade da rede. Dessa forma, busca-se demonstrar, de forma prática, como contratos inteligentes podem atuar como uma camada de segurança confiável, transparente e auditável para o provisionamento de dispositivos em ambientes distribuídos. Para atender a esse propósito, o presente estudo visa cumprir os seguintes objetivos:

- Estabelecer a arquitetura geral da solução para embasar a implementação dos contratos inteligentes e consolidar o modelo proposto.
- Definir a estrutura de dados adequada para representar dispositivos IoT nos smart contracts.
- Implementar mecanismos de controle de acesso para restringir o registro de dispositivos a administradores autorizados.
- Validar o funcionamento dos contratos por meio de cenários de registro, autenticação e restrições administrativas.

Metodologia

Um contrato inteligente minimamente viável (MVSC) é a versão mais simples possível de um contrato inteligente que cumpre seu propósito central, permitindo que os desenvolvedores testem e validem sua funcionalidade e obtenham feedback inicial dos usuários antes de investir em um desenvolvimento em larga escala. Para modelar o MVSC proposto, adotou-se como recorte exemplificativo um conjunto de quatro tipos de sensores — temperatura, proximidade, movimento e umidade — representando um cenário de registro e autenticação de dispositivos IoT. Considerando a diversidade funcional desses dispositivos e visando uma gestão mais modular, o sistema foi organizado em quatro contratos independentes, cada um responsável pela gestão de uma categoria específica de sensores.

Com essa configuração estabelecida, a prototipagem dos contratos inteligentes seguiu um fluxo estruturado em etapas, adotado para viabilizar a evolução incremental das funcionalidades e facilitar a verificação progressiva dos requisitos definidos. Inicialmente, procedeu-se à definição da estrutura de dados necessária à representação dos dispositivos, bem como à elaboração de uma versão preliminar do contrato, contemplando funções fundamentais de registro e autenticação. Posteriormente, foram incorporados os mecanismos de controle de acesso e delineados os procedimentos de implantação e de verificação preliminar do comportamento do modelo. Por fim, a etapa conclusiva abrangeu a elaboração de testes automatizados e a preparação de uma demonstração prática de uso, objetivando assegurar a validação das funcionalidades básicas previstas para os contratos.

Para efeito de simulação, a arquitetura dos *smart contracts* foi concebida de modo a possibilitar a integração com uma *Application Programming Interface* (API), responsável por intermediar chamadas externas — como registros e autenticações — e organizar as requisições de forma controlada. Esse desenho favorece a reutilização do modelo em diferentes contextos, possibilitando a integração dos contratos a aplicações externas e, assim, a realização de simulações mais próximas dos ambientes de produção. Dessa forma, assegura-se não apenas a compatibilidade com os mecanismos definidos nesta pesquisa, mas também a ampliação de seu potencial de aplicação em trabalhos correlatos. A Figura 1

apresenta, de forma esquemática, a arquitetura resultante, destacando os fluxos de interação entre dispositivos IoT, a API e os contratos inteligentes, conforme delineado.

A implementação da arquitetura foi conduzida no ecossistema Ethereum, uma plataforma de blockchain pública e de código aberto que oferece suporte a contratos inteligentes de forma descentralizada. Nesse contexto, os *smart contracts* foram desenvolvidos na linguagem de programação Solidity e, inicialmente, prototipados na IDE Remix, o que possibilitou a verificação rápida da lógica e das funcionalidades básicas previstas. Posteriormente, a validação ocorreu por meio do framework Hardhat em conjunto com a rede local Ganache, permitindo testes mais abrangentes e controlados em ambiente de simulação. Essa etapa incluiu a preparação de scripts de implantação configurados para redes de teste locais, viabilizando a reprodução de interações com a blockchain e a análise do comportamento esperado dos contratos sem a necessidade de conexão a uma rede pública. Entre os cenários avaliados, destacaram-se registros válidos, tentativas de acesso não autorizado, prevenção de duplicidades e autenticação de dispositivos com prazo de validade expirado.

Resultados e Discussões

Os resultados da implementação da arquitetura proposta materializaram-se em quatro contratos inteligentes — *HumiditySensorManager*, *TemperatureSensorManager*, *ProximitySensorManager* e *MotionSensorManager*. A estrutura dos contratos segue um modelo uniforme e é composta por três elementos centrais: a definição de um administrador, que concentra as permissões de registro; uma estrutura (*struct*) para sintetizar os atributos do dispositivo (endereço MAC, tipo de medida, data de registro, prazo de validade e status de validade) de forma a assegurar a unicidade, permitir a identificação física, definir o período de confiança e viabilizar a verificação imediata da situação de cada ativo registrado; e um mapeamento que usa o identificador único (UID) como chave primária para indexação e recuperação das entidades. Complementarmente, eventos foram definidos para registrar, em log, as operações de cadastro realizadas na rede.

A função de registro foi implementada com um mecanismo de prevenção de duplicidade, de modo a rejeitar tentativas de cadastrar um dispositivo cujo UID já conste na estrutura. Adicionalmente, o prazo de validade é automaticamente atribuído no momento do cadastro, estabelecendo uma janela temporal de confiança. A função de autenticação, por sua vez, realiza, em uma única operação, a verificação da existência do dispositivo, da validade lógica de seu estado e da vigência do prazo estipulado. Assim, apenas dispositivos ativos e dentro de sua janela temporal são reconhecidos como legítimos pela rede. Ambos os processos ocorrem diretamente na blockchain, garantindo sua condução descentralizada e reforçando a transparência e a confiabilidade da arquitetura.

Para ilustrar a implementação delineada, o código do contrato *HumiditySensorManager* é apresentado na Listagem 1. Os demais contratos seguem uma estrutura similar, com a única variação no valor atribuído ao atributo *measurementType*, que corresponde ao tipo de sensor.

Com o objetivo de validar o comportamento dos *smart contracts* desenvolvidos, foram elaborados e executados testes automatizados no framework Hardhat, utilizando a biblioteca Chai para asserções. Esses testes tiveram como propósito verificar a correta execução das operações de registro e autenticação dos dispositivos, bem como a aplicação das restrições administrativas previstas na lógica dos contratos. A Tabela 1 sintetiza os cenários avaliados

nos quatro contratos implementados (*HumiditySensorManager*, *TemperatureSensorManager*, *ProximitySensorManager* e *MotionSensorManager*), incluindo os testes de registro de sensores válidos, tentativa de registro duplicado, tentativa de registro com parâmetros vazios, autenticação de dispositivos dentro do prazo de validade, autenticação de dispositivos expirados e tentativa de registro de dispositivos por contas não-admin.

Os resultados demonstram que todos os contratos atenderam de forma consistente aos requisitos definidos. As operações de registro e autenticação foram corretamente processadas, e, em particular, o tratamento de situações de duplicidade e expiração. Além disso, a aplicação da restrição de acesso administrativo assegurou que apenas o endereço responsável pela implantação do contrato pudesse registrar novos sensores.

Conclusão

Este trabalho apresentou o desenvolvimento de uma arquitetura descentralizada para gerenciamento e controle de MVSC para registro e autenticação de dispositivos IoT em uma rede blockchain. A proposta abordou os desafios recorrentes de segurança e confiança em ambientes distribuídos — destacadamente a necessidade de mecanismos descentralizados, transparentes e auditáveis — e demonstrou, na prática, como os *smart contracts* podem atuar como uma camada de segurança para o provisionamento de dispositivos. O escopo contemplou quatro categorias de sensores (temperatura, proximidade, movimento e umidade) e a organização modular em quatro contratos independentes de gerenciamento, concebidos para favorecer a evolução e a validação incremental das funcionalidades previstas.

Do ponto de vista metodológico, a solução foi estruturada a partir de um modelo de arquitetura que prevê a integração com uma API para intermediar chamadas externas (registro e autenticação) e organizar requisições de forma controlada. A implementação ocorreu no ecossistema Ethereum, com codificação em Solidity, prototipagem inicial na IDE Remix e validação posterior no framework Hardhat conectada à rede local Ganache, permitindo simulações reproduzíveis de cenários de uso. Esse arranjo viabilizou a inspeção do comportamento dos contratos em situações de registros válidos, tentativas de acesso não autorizado, prevenção de duplicidade e autenticação após expiração de prazo.

A lógica aplicada nos contratos desenvolvidos consolidou mecanismos de controle de acesso, verificação temporal, prevenção de duplicidade e autenticação confiável, assegurando que apenas dispositivos válidos fossem reconhecidos. Os testes executados confirmaram esses aspectos, demonstrando que a arquitetura proposta e o código implementado foram capazes de atender integralmente aos objetivos definidos, validando o modelo com testes de registro de sensores válidos, tentativa de registro duplicado, tentativa de registro com parâmetros vazios, autenticação de dispositivos dentro do prazo de validade, autenticação de dispositivos expirados e tentativa de registro de dispositivos por contas não-admin.

Dessa forma, o presente estudo consolida um alicerce funcional para o uso de contratos inteligentes na autenticação de dispositivos IoT, demonstrando a viabilidade da integração entre blockchain e Internet das Coisas em cenários descentralizados de registro e verificação de identidade. Como perspectiva futura, destaca-se o aprimoramento dos contratos com a incorporação de funcionalidades de segurança mais complexas, como renovação e revogação de credenciais, controle de acesso baseado em papéis e integração com mecanismos de identidade digital. Assim, a solução aqui apresentada, embora concebida como uma proposta

básica, abre caminho para a evolução em direção a um modelo mais completo, robusto e adequado a aplicações em larga escala.

Referências

- Akkaoui, R. (2021). Blockchain for the management of internet of things devices in the medical industry. *IEEE Transactions on Engineering Management*, 70(8):2707–2718.
- Bargayary, B. and Medhi, N. (2023). A blockchain-assisted authentication for sdn-iot network using smart contract. In *Proceedings of the 4th International Conference on Computing and Communication Systems (I3CS)*.
- Garba, A., Khoury, D. J., Balian, P., Haddad, S., Sayah, J., Chen, Z., Guan, Z., Ham-dan, H., Charafeddine, J., and Mutib, K. A. (2023). Lightcert4iots: Blockchain-based lightweight certificates authentication for IoT applications. *IEEE Access*, 11:28370–28383.
- Goswami, B. and Choudhury, H. (2022). A blockchain-based authentication scheme for 5g-enabled IoT. *Journal of Network and Systems Management*, 30(4):61–85.
- Hameed, K., Garg, S., Amin, M. B., and Kang, B. (2021). A formally verified blockchain-based decentralised authentication scheme for the internet of things. *The Journal of Supercomputing*, 77(7):14461–14501.
- Hussein, D. H., Houlahan, E., Janelle-Goode, A., Lumsden, G., and Ibnekahla, M. (2024). A blockchain-based dual identity management and authentication framework for IoT networks. In *Proceedings of the 2024 IEEE 10th World Forum on Internet of Things (WF-IoT)*, pages 544–549.
- Khashan, O. A. and Khafajah, N. M. (2023). Efficient hybrid centralized and blockchain-based authentication architecture for heterogeneous IoT systems. *Journal of King Saud University — Computer and Information Sciences*, 35(2):726–739.
- Liou, W.-C. and Lin, T. (2021). T-auth: A novel authentication mechanism for the IoT based on smart contracts and pufs. In *2021 IEEE International Conference on Communications Workshops (ICC Workshops)*, pages 1–6.
- Liu, Y., Liu, A., Xia, Y., Hu, B., Liu, J., Wu, Q., and Tiwari, P. (2024). A blockchain-based cross-domain authentication management system for iot devices. *IEEE Transactions on Network Science and Engineering*, 11(1):115–127.
- Mahmoud, C., Aouag, S., Cherroun, H., Brik, B., and Rezgui, A. (2024). A decentralized blockchain-based authentication scheme for cross-communication in IoT networks. *Cluster Computing*, 27(3):2505–2523.
- Majd, N. E. and Sharko, M. (2023). Secure and cost effective IoT authentication and data storage framework using blockchain nft. In *Proceedings of the 32nd International Conference on Computer Communications and Networks (ICCCN)*, pages 1–7.
- Qiu, S., Li, J., Di, X., Li, X., Wu, Y., and Ibrahim, M. (2024). Lightweight mutual authentication scheme based on blockchain for the internet of medical things. *IEEE Internet of Things Journal*, 12(7):8848–8861.

Anexos

Figura 1 - Arquitetura de Gerenciamento e controle de MVSC

Listagem 1 – Estrutura do contrato HumiditySensorManager

Tabela 1 - Resultados dos testes automatizados para os contratos de provisionamento de sensores

CÓDIGO: ET0414

AUTOR: NICOLE DANTAS LOPES

COAUTOR: JOAO PEDRO BARRETO MELLO

ORIENTADOR: BRUNO SANTANA DA SILVA

TÍTULO: Uma Avaliação de Comunicabilidade do BioTax

Resumo

A classificação dos seres vivos com uma taxonomia auxilia a organização do conhecimento e a compreensão da enorme biodiversidade. A identificação taxonômica de espécimes é uma atividade muito comum no estudo, na prática profissional e na pesquisa em Biociências. Pela grande quantidade de informações envolvidas, um software pode auxiliar na identificação taxonômica. Um desses softwares é o BioTax, que vem sendo desenvolvido na UFRN.

Este trabalho teve por objetivo avaliar a comunicabilidade da interface do BioTax para verificar se os usuários enfrentam dificuldades durante seu uso. Cinco estudantes de graduação em Biologia da UFRN participaram deste estudo. Eles foram convidados a utilizar a versão 0.0.1 desktop do software para realizar oito tarefas. Sua interação foi observada e registrada para análise de acordo com o Método de Avaliação de Comunicabilidade.

Os participantes enfrentaram problemas de comunicabilidade na maior parte das tarefas, desde problemas mais simples e temporários até dificuldades que os impediram de concluir a tarefa. As tarefas de identificação de espécimes e de cadastro de chave de identificação foram as mais afetadas por problemas de comunicabilidade. Trabalhos futuros precisam investigar melhor a causa desses problemas de comunicabilidade e encaminhar melhorias na interface do BioTax.

Palavras-chave: Comunicabilidade. Interface. Biologia. Taxonomia. Chave de identificação

TITLE: A Communicability Evaluation of BioTax

Abstract

Classifying living beings using a taxonomy helps organize knowledge and understand the vast biodiversity. Taxonomic identification of specimens is a very common activity in studies, professional practice, and research in the Biosciences. Due to the large amount of

information involved, software can assist in taxonomic identification. One such software is BioTax, which is being developed at UFRN.

This study aimed to evaluate the communicability of the BioTax interface to determine whether users encounter difficulties using it. Five undergraduate Biology students from UFRN participated in this study. They were asked to use the software's desktop version 0.0.1 to perform eight tasks. Their interactions were observed and recorded for analysis according to the Communicability Assessment Method.

Participants faced communicability issues in most tasks, ranging from simple and temporary issues to difficulties that prevented them from completing the task. Specimen identification and identification key registration tasks were the most affected by communicability issues. Future work needs to further investigate the cause of these communication problems and address improvements to the BioTax interface.

Keywords: Communicability. Interface. Biology. Taxonomy. Identification key.

Introdução

A taxonomia dos seres vivos (Linnaeus, 1758) auxilia a organizar sua enorme diversidade em grupos hierárquicos (Amorin, 2002). Ela permite estruturar, facilitar o acesso e até raciocinar sobre o conhecimento biológico (Rosa; Martins, 2007). Por exemplo, a taxonomia é muito útil para ajudar no entendimento das relações evolutivas e ecológicas dos seres vivos (Brasil, 2018). Então, a identificação taxonômica de espécimes faz parte do estudo, da prática profissional e da pesquisa científica em muitas áreas da Biociência.

A identificação taxonômica muitas vezes é orientada por um instrumento chamado de chave de identificação (Walter; Winterton, 2007). Uma chave de identificação lembra uma lista de verificação que precisa ser conferida para se determinar que um certo espécime (um exemplar de ser vivo) pertence a um determinado táxon, ou seja, a determinados grupos hierárquicos na taxonomia. A chave de identificação pode ser dicotômica, com apenas dois conjuntos de características a serem verificadas por vez (Oliveira-Costa, 2013), ou politômica, que oferecem muitas opções de verificação a cada momento (Rocha et al., 2019).

Existem muitos softwares para apoiar a identificação taxonômica (Dallwitz, 2020; Walter; Winterton, 2007). Um deles é o BioTax (Figuras 1, 2 e 3). Este software desktop tem sido desenvolvido na UFRN (Silva et al., 2021) para permitir que o usuário realize identificações taxonômicas com as chaves dicotômicas que desejar. Estas chaves podem ser criadas por outras pessoas e importadas para uso, ou criadas no próprio BioTax conforme o usuário considerar adequado.

O BioTax tem sido utilizado como um material didático em cursos superiores de Biociências (Gama et al., 2022). Entretanto, pouco tem sido investigado sobre a qualidade de uso da sua interface com usuário. A qualidade de uso possui quatro critérios definidos na área de Interação Humano-Computador: usabilidade, comunicabilidade, experiência do usuário e acessibilidade (Barbosa; Silva, 2010). Em trabalhos anteriores, a usabilidade do BioTax já foi avaliada. Neste trabalho, o foco da avaliação passou a ser a comunicabilidade. A

comunicabilidade indica a capacidade da interface comunicar de forma eficiente e eficaz para o usuário qual é a lógica de design (Prates et al., 2020).

Então, o objetivo deste trabalho foi avaliar a comunicabilidade do BioTax com a participação de estudantes de Biociências no Ensino Superior.

Metodologia

O método empregado para avaliar a interface do BioTax foi o Método de Avaliação de Comunicabilidade (MAC) (Prates et al., 2020). Este é um método de avaliação que observa e analisa o uso que os usuários fazem do sistema avaliado (Barbosa; Silva, 2010). Cada participante da avaliação foi convidado a usar o BioTax para realizar algumas tarefas. Sua interação foi observada, registrada para posterior análise conforme proposto pelo MAC.

A coleta de dados desta avaliação ocorreu no Laboratório de Insetos e Vetores (LIVE) da Universidade Federal do Rio Grande do Norte (UFRN). O uso que cada participante fez individualmente do BioTax foi observado. O vídeo da tela do computador e o áudio do que estava sendo falado foram gravados com o software Free Cam 8.

Antes das tarefas serem apresentadas ao participante, foi solicitado que ele respondesse um formulário online sobre seu perfil, disponibilizado via Google Forms. As perguntas deste questionário abordaram o uso de tecnologias digitais no dia a dia dos participantes e suas experiências com chaves de identificação, tanto em papel quanto digitais.

O BioTax foi disponibilizado ao participante em um computador portátil. Depois de terem a oportunidade de fazer uma exploração livre do BioTax, os participantes receberam a demanda para a executar as seguintes tarefas:

- Tarefa 1 - Importar chave e ler a chave
- Tarefa 2 - Identificar a taxonomia de 3 exemplares
- Tarefa 3 - Identificar a taxonomia de 1 exemplar
- Tarefa 4 - Exportar os resultados da identificação
- Tarefa 5 - Cadastrar táxons
- Tarefa 6 - Cadastrar chave
- Tarefa 7a - Remover figura da chave
- Tarefa 7b - Substituir figura na chave
- Tarefa 7c - Remover passo da chave
- Tarefa 8 - Exportar a chave

Além do computador com BioTax, para realizar as Tarefas 1, 2, 3 e 4, os participantes também receberam: uma chave de identificação de moscas, outra chave de identificação de percevejos, três espécimes de moscas, um espécime de percevejo, além de lupas comum e eletrônica. Para a criação de uma chave de identificação, Tarefas 5 até 8, os participantes

receberam a taxonomia necessária a ser cadastrada, um documento com a descrição da chave de interesse e arquivos de imagem necessários.

Ao fim das tarefas foi realizada uma breve entrevista com cada participante com perguntas sobre sua experiência de uso do software.

Na análise intrasujeito, o vídeo gravado com cada participante foi analisado para identificar as rupturas de comunicação que ocorreram, considerando as 13 etiquetas proposta pelo MAC:

- Cadê?
- E agora?
- O que é isto?
- Epa!
- Onde estou?
- Ué, o que houve?
- Por que não funciona?
- Assim não dá.
- Vai de outro jeito.
- Não, obrigado!
- Pra mim está bom.
- Socorro!
- Desisto.

Depois de concluir a análise da interação de todos os participantes individualmente, realizou-se a análise intersujeito com uma interpretação conjunta das rupturas de comunicação identificadas para tentar identificar oportunidades de melhoria na interface do BioTax. Considerou-se principalmente o contexto, a sequência e a frequência de ocorrência das etiquetas das rupturas de comunicação.

Resultados e Discussões

Os participantes do estudo foram cinco alunos do 5º semestre do curso de Bacharelado em Biologia da Universidade Federal do Rio Grande do Norte (UFRN). Todos os participantes utilizam celulares várias vezes por dia (Figura 4). Quanto ao uso de computadores ou notebooks, 60% dos participantes os utilizam várias vezes por dia e 40% os utilizam apenas uma vez por dia. Todos os participantes utilizam esses dispositivos para assistir vídeos, ouvir música ou áudio, ler notícias e semelhantes e enviar e receber e-mails. A maioria também usa esses dispositivos para conversar com outras pessoas e jogar, enquanto uma minoria também os utiliza para estudar e realizar trabalhos acadêmicos.

Em relação ao uso de chaves de identificação, todos os participantes relataram ter utilizado chaves de identificação em papel durante suas disciplinas, como "Parasitologia para o Ensino de Ciências", "Entomologia Geral II", "Metazoa II" e "Sistemática de Angiospermas". Das chaves utilizadas, todos já utilizaram chaves em papel somente com texto, enquanto uma minoria também utilizou com desenho e apenas um participante com foto real. A experiência com essas chaves em papel foi variada, sendo classificada como regular pela maioria, enquanto uma avaliação foi negativa e outra positiva. De todos participantes, apenas um já utilizou chaves de identificação digitais anteriormente e achou a experiência regular.

Sobre as interações observadas, a Tabela 1 apresenta a quantidade de ocorrências das rupturas de comunicação durante a interação de cada participante ao longo da realização de todas as tarefas solicitadas. Os Participantes 2, 3 e 5 enfrentaram rupturas de comunicação completas com resultados finais diferentes do esperado nas tarefas. Cada um concluiu uma atividade sem perceber que obtiveram um resultado distinto do solicitado ("Para mim está bom"). O Participante 2 desistiu duas vezes de continuar tentando realizar as tarefas, mesmo sem ainda ter chegado no resultado proposto. Esses resultados são os mais graves pois indicam que os participantes não foram capazes de concluir as tarefas por problemas de comunicabilidade.

Todos os participantes enfrentaram falhas de comunicação parciais, que os impediram de obter uma parcela relevante dos resultados esperados. Nesses casos, eles não foram capazes de entender a interação e interface proposta para a execução das tarefas ("Vai de outro jeito"). Este tipo de dificuldade ocorreu 3 vezes para os Participantes 1, 3 e 5, e quatro vezes para os Participantes 2 e 4. Essa falta de compreensão da solução de IHC foi bem frequente no BioTax para todos os participantes e precisa ser tratada adequadamente.

Todos os participantes precisaram interromper temporariamente o fluxo de interação dedicado à realização da tarefa para se recuperar de alguma ruptura de comunicação. Nesse sentido, os Participantes 2, 3, 4 e 5 tiveram dificuldades de encontrar uma forma de expressar suas intenções na interface do BioTax ("Cadê?"). Em uma situação, o Participante 2 não conseguiu perceber ou entender o que estava apresentado na interface ("Ué, o que houve?"). Todos os participantes passaram por alguma situação em que não foram capazes de formular sua próxima intenção de comunicação durante a interação ("E agora?").

Novamente, todos os participantes passaram por situações onde perceberam que não foram compreendidos pelo sistema do modo esperado ao se expressarem na interface. Os Participantes 1, 2, 3 e 4 perceberam que expressaram várias vezes algo num contexto errado ("Onde estou?"). Todos os participantes, em especial o Participante 2, percebeu que cometeu um equívoco ao falar algo não apropriado. Apesar deste equívoco do usuário nem sempre ter relação com a solução de IHC, a interação ou a interface podem estar aumentando o risco de equívocos ou deixando de evitá-los. Os Participantes 2, 4 e 5 perceberam que não seria possível obter os resultados esperados depois de se engajarem em determinados processos de interação ("Assim não dá."). Eles tiveram que abandonar essa estratégia original para tentar concluir as tarefas propostas.

Todos os participantes tiveram dificuldades para compreender o que foi comunicado pelo sistema (preposto designer) através da interface. Isso ocorreu quando o usuário buscou sob demanda acessar mais partes da metacomunicação dentro da própria interface para entender o que foi comunicado ("O que é isso?"). Os Participantes 2, 3 e 5 também precisaram consultar sob demanda mais partes da metacomunicação fora da própria interface para entender o que foi comunicado ("Socorro!"). O Participante 2 enfrentou uma situação em que teve dificuldade para compreender porque o sistema comunicou um resultado diferente do

esperado (“Por que não funciona?”). Nesses casos, o Participante 2 se destacou dos demais pelo número de vezes que enfrentou problemas na compreensão do que o sistema comunicou.

A Tabela 2 apresenta a frequência de ocorrência de problemas de comunicação observados nesta avaliação. A Tarefa 6 de cadastrar uma chave de identificação foi a que mais acumulou problemas (31) de comunicabilidade, com destaque para as ocorrências das etiquetas “Epa!” (9 vezes) e “O que é isso?” (8 vezes). Parte dessas dificuldades são compreensíveis e até aceitáveis, pois esta é uma tarefa muito complexa que os usuários não estão acostumados a realizar tanto no meio digital quanto analógico. A Tarefa 3, de identificar a taxonomia do segundo grupo de espécimes fornecido, também acumulou muitos problemas de comunicabilidade (15). Apesar desta ser uma repetição da tarefa anterior com um outro grupo de espécimes, é importante lembrar que os participantes foram novamente estimulados a explorar na interface um suporte específico desenvolvido no BioTax. Ficou evidente que a tentativa de usar esta nova funcionalidade gerou muitos problemas de comunicabilidade. A Tarefa 5, de cadastrar táxons, também contabilizou muitos problemas de comunicabilidade (15). Ainda que o conjunto de táxons necessários tenha sido fornecido por escrito, com uma representação visual indentada para indicar a hierarquia, os participantes enfrentaram muitas dificuldades. Parte dessas dificuldades pode estar relacionada com a falta de hábito dos participantes de realizar tarefas que definem táxons dentro de uma taxonomia. Então, surge uma questão, essa tarefa de cadastro de táxon realmente é necessária se ela não é tão comum no cotidiano dos usuários?

Os problemas de comunicabilidade mais graves ocorreram nas Tarefas 1 e 6, de importação e de cadastro de chave, respectivamente. Os problemas de comunicabilidade medianos ocorreram nas Tarefas 2, 3 e 6, de identificação taxonômica e de cadastro de parte de uma taxonomia.

A necessidade de trocar o bolsista no meio das atividades e os problemas pessoais dos envolvidos em diferentes momentos impediram um aprofundamento da análise desses resultados.

Conclusão

A avaliação de comunicabilidade da interface do BioTax identificou que os usuários enfrentam muitas dificuldades durante o uso por problemas na sua comunicação com o sistema. Isso prejudicou a interação com o BioTax durante a realização de tarefas típicas previstas para o suporte do sistema. Deste modo, é importante que a interface com usuário do BioTax seja aprimorada para que os usuários possam usufruir melhor dos recursos tecnológicos oferecidos por este software. Trabalhos futuros deveriam realizar mais avaliações da interface do BioTax e investigar formas de melhorar a sua interface com usuário.

Referências

- AMORIM, D. S. *Fundamentos de sistemática filogenética*. Holos, 2002.
- BARBOSA, S. D. J.; SILVA, B. S. *Interação Humano-Computador*. Elsevier, 2010.
- BRASIL, Ministério da Educação. Base Nacional Comum Curricular - BNCC. 2018.

DALLWITZ, M.J. Programs for interactive identification and information retrieval. 2020. Disponível em: <https://www.delta-intkey.com/www/idprogs.htm>. Acesso em: 26 ago. 2025.

GAMA, R. A.; BARBOSA, T. M.; SILVA, B. S. Sequência didática com uso de chave de identificação digital em disciplina remota de Entomologia Médica na graduação. **Revista de Ensino de Ciências e Matemática**, v. 13, n. 2, p. 1–25, 2022.

LINNAEUS, C. V. *Systema Naturae per regna tria naturae*. 10^a ed., v. 1, p. 823, 1758.

OLIVEIRA-COSTA, J. *Insetos peritos: a entomologia forense no Brasil*. Campinas: Millenium, 2013.

PRATES, R.O.; DE SOUZA, C.S.; BARBOSA, S.D.J. A Method for Evaluating the Communicability of User Interfaces. In *ACM Interactions*, Jan-Feb, pp 31-38, 2000.

ROCHA, D. A. et al. LutzDex™—A digital key for Brazilian sand flies (Diptera, Phlebotominae) within an Android App. *Zootaxa*, v. 4688, n. 3, p. 382–388, 2019.

ROSA, J. M.; MARTINS, L. M. Reflexões sobre o ensino da taxonomia e da sistemática filogenética e o desenvolvimento do pensamento abstrato. Obutchénie: **Revista de Didática e Psicologia Pedagógica**, v. 1, n. 2, p. 376–410, 2017.

SILVA, B. S.; LIMA, P. G. S.; GAMA, R. A. **BioTax Desktop**. 2021.
Patente: Programa de Computador. Número do registro: BR512021002690-0.

WALTER, D.E.; WINTERTON, S. Keys and the crisis in taxonomy: extinction or reinvention? **Annu. Rev. Entomol.**, v. 52, p. 193–208, 2007.

Anexos

Figura 1 - Chaves importadas no BioTax

Figura 2 - Consulta da chave no BioTax

Figura 3 - Taxonomia cadastrada no BioTax

Figura 4 - Perfil dos participantes

Tabela 1 - Quantidade de problemas de comunicação por participante

Tabela 2 - Quantidade de problemas de comunicação por tarefa

CÓDIGO: ET0422

AUTOR: RODRIGO FERREIRA DA NOBREGA

COAUTOR: ANDERSON CLAUDIO RODRIGUES DA SILVA

COAUTOR: SAULO CARLOS PEREIRA DA ROCHA

ORIENTADOR: ROGER KREUTZ IMMICH

TÍTULO: Utilização de Smart Contracts para segurança de dispositivos IoT na Rede Blockchain da Ethereum

Resumo

Este artigo explora a viabilidade da utilização da tecnologia blockchain na rede Ethereum em conjunto com Smart Contracts, para aprimorar a segurança no processo de registro de dispositivos de Internet das Coisas (IoT). A arquitetura utiliza uma API RESTful como intermediário entre os dispositivos IoT e a blockchain, facilitando a integração e abstraindo a complexidade da rede para os clientes. O desenvolvimento dos Smart Contracts foi realizado em Solidity, com uma abordagem modularizada para cada tipo de sensor, visando avaliar o impacto da modularização na autenticação. O backend foi desenvolvido em Java 21 com Spring Boot 3 e containerizado com Docker. Testes de estresse foram conduzidos utilizando Apache JMeter em cenários isolados e concorrentes para avaliar o desempenho do sistema sob diferentes cargas. Os resultados revelaram que, embora a arquitetura seja funcional e garanta a integridade transacional, a serialização das transações para evitar problemas de nonce introduz um gargalo de latência. A capacidade de processamento do nó Ganache foi identificada como um fator limitante para a escalabilidade em cenários de alta concorrência, evidenciando um trade-off entre velocidade de processamento e garantia da integridade na blockchain. Conclui-se que a solução proposta é eficaz para a segurança do registro de dispositivos IoT, mas sua escalabilidade em ambientes de alta demanda requer otimizações na camada blockchain.

Palavras-chave: Blockchain, Contratos Inteligentes, IoT, API, Autenticação, Ethereum

TITLE: Using Smart Contracts to Secure IoT Devices on the Ethereum Blockchain Network

Abstract

This paper explores the feasibility of using blockchain technology on the Ethereum network in conjunction with Smart Contracts to improve security in the Internet of Things (IoT) device registration process. The architecture uses a RESTful API as an intermediary between IoT devices and the blockchain, facilitating integration and abstracting network complexity for clients. The Smart Contracts were developed in Solidity, with a modularized approach for each sensor type, aiming to evaluate the impact of modularization on authentication. The backend was developed in Java 21 with Spring Boot 3 and containerized with Docker. Stress tests were conducted using Apache JMeter in isolated and concurrent scenarios to evaluate

system performance under different loads. The results revealed that, although the architecture is functional and ensures transactional integrity, serializing transactions to avoid nonce issues introduces a latency bottleneck. The processing capacity of the Ganache node was identified as a limiting factor for scalability in high-concurrency scenarios, highlighting a trade-off between processing speed and ensuring blockchain integrity. It is concluded that the proposed solution is effective for the security of IoT device registration, but its scalability in high-demand environments requires optimizations at the blockchain layer.

Keywords: Blockchain, Smart Contracts, IoT, API, Authentication, Ethereum

Introdução

O aumento crescente de dispositivos de Internet das Coisas (IoT) tem transformado diversos setores, desde a automação residencial e indústria 4.0 (CARRION 2019), proporcionando um alto nível de conectividade e automação. No entanto, essa expansão traz consigo desafios significativos, especialmente no que tange à segurança e à integridade dos dados (LEITE 2019). A natureza distribuída e, muitas vezes, heterogênea dos dispositivos IoT os torna alvos potenciais para ataques cibernéticos, comprometendo a privacidade dos usuários e a funcionalidade dos sistemas. A segurança do registro e da autenticação desses dispositivos é, portanto, uma preocupação primordial para garantir a confiabilidade e a robustez das redes IoT.

Dessa forma, a tecnologia blockchain surge como uma solução promissora para endereçar as vulnerabilidades inerentes aos sistemas centralizados tradicionais. A blockchain, com sua arquitetura descentralizada, imutabilidade e mecanismos criptográficos, oferece um ecossistema robusto para a criação de registros seguros e auditáveis (RIBEIRO 2021). A aplicação de Smart Contracts, programas auto executáveis armazenados na blockchain, permite automatizar e impor regras de segurança de forma transparente e à prova de adulteração (BUTERIN 2013).

Este artigo investiga a aplicação de Smart Contracts na rede Ethereum para fortalecer a segurança do processo de registro de dispositivos IoT. O objetivo principal é demonstrar a viabilidade e os benefícios de integrar a tecnologia blockchain em uma arquitetura de segurança para IoT, abordando as limitações dos modelos atuais e propondo uma solução inovadora. A pesquisa detalha o desenvolvimento de um sistema que utiliza uma API RESTful como ponte entre os dispositivos IoT e a blockchain, facilitando a interação e garantindo que as operações de registro sejam seguras e verificáveis.

Serão apresentados os detalhes da metodologia empregada, que abrange o desenho da arquitetura do sistema e a realização de testes de desempenho rigorosos. A escolha da linguagem Solidity para os Smart Contracts e a abordagem modularizada para o registro de diferentes tipos de sensores serão discutidas, juntamente com a infraestrutura de backend desenvolvida em Java com Spring Boot. A validação da solução proposta foi realizada por meio de testes de estresse utilizando Apache JMeter, simulando cenários de carga isolados e concorrentes para avaliar a latência, o throughput e a escalabilidade do sistema.

Os resultados e a discussão analisaram o desempenho da arquitetura sob diferentes condições de estresse, identificando gargalos e trade-offs entre a integridade transacional e a velocidade de processamento. Em particular, será examinada a questão da serialização de transações para evitar problemas de nonce e seu impacto na escalabilidade. Este estudo visa contribuir para o avanço das soluções de segurança para IoT, oferecendo insights práticos sobre a integração

da blockchain e Smart Contracts, e delineando as futuras direções para otimizar o desempenho em ambientes de alta demanda.

Metodologia

A presente pesquisa foi conduzida para demonstrar a viabilidade da utilização da tecnologia blockchain para garantir a segurança no processo de registro dos dispositivos de Internet das Coisas (IoT). A metodologia foi dividida em três fases principais: o desenho da arquitetura do sistema, a implementação dos componentes de software e a realização de testes funcionais para validar a solução proposta.

A Figura 1 mostra o diagrama da arquitetura desenvolvida, exibindo as três camadas do sistema. A camada mais externa é a Camada de Cliente, composta pelos dispositivos IoT. A camada intermediária é a Camada de Serviço, implementada como uma API RESTful, que atua como um intermediário entre o cliente e a blockchain. Por fim, a camada base é a Camada de Persistência e Confiança, materializada pela rede blockchain.

Nesse modelo, os clientes seriam os dispositivos que irão realizar requisições para a API utilizando métodos HTTP como POST ou GET, enviando apenas dois dados, o UUID do dispositivo e o endereço MAC. No modelo proposto cada endpoint é responsável por redirecionar a requisição para alguma funcionalidade específica dos contratos inteligentes. Dessa forma, a API foi desenvolvida para facilitar a integração entre os dispositivos e a rede blockchain, uma vez que a camada dos clientes não precisam conhecer o endereço dos Smart Contracts, como também da rede blockchain.

A blockchain exerce o papel fundamental no desenvolvimento desse sistema, ela atua como uma camada de persistência e confiança, garantindo todas as características necessárias para a auditabilidade do processo de autenticação. No desenvolvimento dos contratos, optou-se por uma abordagem priorizando a segmentação de cada funcionalidade, ou seja, foi desenvolvido um contrato específico para registrar cada tipo específico de sensor.

Para o desenvolvimento do serviço de backend, a seleção tecnológica foi pautada em critérios de robustez, portabilidade e manutenibilidade. Visando a construção de uma aplicação escalável e eficiente.

A linguagem de programação Java 21 foi escolhida como base devido à sua arquitetura multiplataforma. Além disso, a maturidade da plataforma Java, seu ecossistema de bibliotecas e o forte suporte à programação concorrente foram considerados fatores determinantes para a construção de um serviço escalável e resiliente, capaz de lidar com altas cargas de processamento.

Sobre essa base, foi empregado o framework Spring Boot 3, que agiliza o processo de desenvolvimento e configuração de aplicações backend. A adoção do Spring se fundamenta em sua arquitetura baseada nos princípios de Inversão de Controle (IoC) e Injeção de Dependência (DI), que fomentam um design de software de baixo acoplamento e facilita a manutenção e a evolução do sistema. A automação do ciclo de vida do projeto e o gerenciamento de dependências foram realizados com o Apache Maven 3.9.10.

Por fim, para garantir um ambiente de desenvolvimento e implantação padronizado e isolado, a aplicação foi integralmente containerizada com o uso da plataforma Docker 28.3.3.

Para o desenvolvimento dos Smart Contracts foi escolhida a linguagem Solidity 0.8.20 devido a sua maturidade e consistência no ambiente blockchain e para simular a rede blockchain foi utilizado o Ganache na versão 2.7.1. O padrão de arquitetura para os contratos seguiu uma abordagem modularizada, ou seja, foi desenvolvido um mini-contrato para registrar cada tipo de sensor específico. Esse método de desenvolvimento foi escolhido para que seja possível avaliar o impacto que a modularização pode causar no processo de autenticação.

A integração entre a camada de serviço Java e os contratos inteligentes foi alcançada por meio da geração de classes wrapper type-safe a partir das Application Binary Interfaces (ABIs) dos contratos. Para essa finalidade, foi empregada a ferramenta de linha de comando web3j, na versão 1.7.0. Durante este processo, foi identificada uma incompatibilidade que resultava em erros de importação nas classes geradas. Como contorno, foi necessário realizar uma modificação manual nos wrappers, consistindo em comentar um bloco de código específico — referente a funcionalidades não utilizadas no escopo deste projeto — para garantir a compilação e a execução bem-sucedida da aplicação.

Com o intuito de testar a arquitetura proposta em cenários reais de estresse e processamento elevado foi utilizada a ferramenta Apache JMeter 5.6.3, os testes foram realizados em três cenários distintos de estresse: o primeiro foi realizado com 250 usuários virtuais, o segundo com 500 usuários virtuais e o terceiro com 1000 usuários virtuais. A quantidade de usuários foi escolhida com a finalidade de promover um estresse significativo para entender como os sistemas lidariam com uma alta quantidade de requisições simultâneas.

Para a avaliação de performance, foram definidos dois cenários de teste distintos. O primeiro, um teste de carga isolado, direcionou requisições de autenticação para um único tipo de sensor por vez, a fim de medir a performance individual de cada contrato. O segundo, um teste de carga concorrente, simulou requisições simultâneas para todos os tipos de sensores, visando avaliar o sistema sob uma carga de trabalho mista e mais próxima da realidade.

Em suma, a metodologia empregada resultou na construção de um sistema funcional e testável, integrando uma API de serviço para redirecionamento com uma camada de persistência descentralizada. A validação do sistema foi conduzida por meio de testes de estresse com a ferramenta Apache JMeter, utilizando cenários de carga progressiva e requisições concorrentes para avaliar a performance e a escalabilidade da arquitetura proposta. Os resultados quantitativos obtidos a partir destes testes são apresentados e analisados na seção subsequente, a fim de determinar a viabilidade da solução.

Resultados e Discussões

Nesta seção, são apresentados e analisados os resultados dos testes de carga realizados sobre a arquitetura da API desenvolvida. O objetivo principal desta avaliação é quantificar o desempenho do sistema sob diferentes níveis de estresse, utilizando indicadores como Latência Média e o Throughput (Vazão).

A análise busca identificar os limites da arquitetura desenvolvida, como também a performance dos contratos inteligentes desenvolvidos de formas separadas para cada sensor.

O ambiente controlado responsável por gerar os testes, foi executado em um ambiente controlado com a seguinte configuração de hardware e software: Arch Linux, Intel® Core™ i5-12500H (12 Cores, 16 Threads) e 16 GB DDR4 3200MHz.

O controle de acesso aos contratos inteligentes foi projetado para ser centralizado, de modo que as operações de escrita, responsáveis por modificar o estado da blockchain, fossem de exclusividade da conta que realizou a implantação. Este modelo de propriedade única foi escolhido para garantir a máxima integridade e controle sobre o processo de registro de novos dispositivos no sistema.

Contudo, uma consequência direta deste design centralizado foi identificada durante os testes de estresse. Ao submeter a API a um alto volume de requisições simultâneas, foram observados erros de validação de nonce. Este comportamento ocorre devido a uma condição de corrida: múltiplos threads da aplicação tentam enviar transações a partir da mesma conta, obtendo o mesmo valor de nonce antes que a transação anterior seja confirmada pela rede, resultando na criação de uma transação com nonce inválido e consequentemente sendo rejeitado pelo processo de validação dos nós da rede blockchain.

Para solucionar este desafio de concorrência, adotou-se uma estratégia de serialização das transações na camada de serviço. Por meio da funcionalidade `synchronized` da linguagem Java, foi criada uma fila virtual que garante o envio das transações de forma sequencial e ordenada, evitando o problema de nonce inválido.

Embora esta solução tenha se mostrado eficaz para garantir que todas as transações fossem processadas sem conflitos, ela introduz um gargalo deliberado no sistema. Portanto, a elevada latência registrada nos resultados dos testes é uma consequência direta e esperada deste mecanismo de sincronização, evidenciando um trade-off fundamental entre a velocidade de processamento e a garantia da integridade transacional na blockchain.

A validação funcional da arquitetura foi realizada por meio de testes individuais para cada endpoint, utilizando a ferramenta de cliente REST Insomnia 11.2.0. A Figura 2 representa a utilização da ferramenta Insomnia para realizar a requisição POST de cadastro do sensor de umidade para a API, mostrando o body enviado e a resposta gerada. A operação de cadastro do sensor de umidade com UUID `ab1c48ae-193d-475a-a08d-8487c1a4e8eb` e Mac Address `43-19-FF-E4-CC-BB` foi concluída com sucesso, conforme evidenciado pela resposta HTTP 200 OK e o body contendo informações necessárias para identificar o tempo de validade do registro.

Complementarmente, a Figura 3 exibe os logs gerados pela camada de serviço da API durante a execução do processo de autenticação do sensor de umidade. Estes registros do backend são apresentados para demonstrar a rastreabilidade interna do sistema, detalhando as etapas do fluxo da requisição na API, desde o recebimento da solicitação até a interação com o contrato inteligente e a resposta da API.

Durante os dois testes de carga foi analisado o consumo de recursos da máquina de teste, e a taxa de uso do processador ficou entre 40% a 60% e o uso da memória ram entre 4Gb a 6Gb, a depender da carga de processamento do teste realizado. Devido a isso, a causa mais provável do gargalo foi devido a camada blockchain Ganache visto que por ser uma rede de teste pode não ter um processamento eficiente das requisições.

Na primeira fase da avaliação, os testes de performance foram conduzidos em um cenário isolado para cada tipo de registro. Nesta abordagem, as requisições de carga foram direcionadas a um único endpoint por vez, evitando a interferência de operações concorrentes com outros contratos ou gasto de processamento da API com outros métodos. O objetivo deste método foi avaliar o desempenho individual de cada smart contract, estabelecendo uma linha de base (baseline) para a latência e a vazão (throughput) de cada operação de forma independente.

A Figura 4 mostra o throughput médio das requisições isoladas e concorrentes dos quatro contratos inteligentes. Ao analisar o desempenho da arquitetura desenvolvida em relação ao throughput é possível verificar que a API possui um limite na capacidade de processamento das requisições. Esse comportamento de estabilidade do throughput pode indicar que ocorreu um limite na capacidade de processamento das requisições/transações.

A Figura 5 informa a latência média dos endpoints ao realizar o teste de forma isolada e concorrente dos quatro contratos inteligentes. Em contraste com a vazão, ao analisar a latência foi identificado um aumento linear de acordo com o aumento da quantidade de usuários fazendo as requisições simultaneamente, aumentando 2,2x de 250 requisições para 500 requisições e 2,1x de 500 requisições para 1000.

O aumento linearmente da latência ocorre devido ao gargalo da taxa de processamento (throughput), ao aumentar a quantidade de usuários simultâneos, consequentemente a quantidade de requisições, isso resulta em uma fila de espera maior, aumentando a latência do sistema.

Na segunda fase da avaliação, os testes de performance foram conduzidos em um cenário concorrente, ou seja, todos os endpoints foram testados simultaneamente. Nesse modelo de teste, o objetivo foi avaliar a performance da arquitetura quando todos os endpoints de registro são requisitados simultaneamente, se aproximando de um ambiente de teste real.

Para esse modelo é importante entender o funcionamento do plano de execução de testes do Apache JMeter. Os usuários virtuais criados executam as requisições em uma ordem sequencial para cada endpoint, a sequência utilizada foi: 1º Umidade, 2º Temperatura, 3º Proximidade e 4º Movimento.

Esse modelo impacta diretamente a validade da comparação de performance entre os endpoints. O endpoint de Umidade beneficia-se por ser o primeiro a ser executado, se beneficiando de um estado de sistema com menor carga. Enquanto o endpoint de Movimento por ser o último a ser executado também conta com uma carga menor de processamento, visto que outros usuários finalizaram o plano de teste dos outros endpoints.

A Tabela 1 exibe de forma condensada os resultados do teste de carga juntamente com o desvio padrão. O comportamento explicado acima fica evidente ao analisar o desvio padrão das variáveis nesse cenário de teste concorrente. É possível perceber que ocorreu uma variabilidade maior nos valores quando comparado aos testes no cenário isolado.

A análise dos resultados dos testes concorrentes demonstrou um impacto significativo no desempenho do sistema quando comparado aos testes isolados. A latência média das requisições aumentou consideravelmente, mas a métrica mais afetada foi o throughput, que

apresentou uma queda de performance de até 50% em determinados cenários de carga, indicando uma menor capacidade de processamento de requisições por segundo.

Ademais, os resultados obtidos levam a crer que a quantidade de 500 requisições foi o valor mais estável encontrado nos resultados apresentados, devido ao desvio padrão menor quando comparado aos outros grupos de requisições, mostrando que pode ser o resultado mais adequado para a simulação realizada.

Esses resultados foram impactados principalmente pela fila de processamento virtual criada para processar o envio das transações à blockchain. Indicando que essa arquitetura não possui um desempenho satisfatório quando submetido a altas cargas de processamento de forma paralela, ou seja, apesar de garantir a correteza das transações a arquitetura não é escalável em cenários de alta concorrência.

A estrutura modular e padronizada utilizada no desenvolvimentos dos Smart Contracts, resultou em métricas de desempenho semelhantes para Latência Média e Throughput. As pequenas variações observadas nos resultados podem ser atribuídas a fatores externos, como a concorrência por recursos de CPU e memória na máquina de teste, gerada por outros processos em execução no sistema.

Um desafio inerente à arquitetura blockchain adotada é a gestão do custo das transações, conhecido como gas. Diferente de sistemas centralizados, toda operação que modifica o estado da rede incorre em um custo que deve ser pago pela conta que criou a transação. Esta característica impõe um requisito operacional de gerenciamento de fundos para a conta de administrador, que atua como o principal agente transacional do sistema.

A criticidade deste requisito foi validada durante a fase de testes. Observou-se que a execução contínua de transações de registro levaria ao esgotamento do saldo da conta administradora na rede Ganache, resultando em falhas de transação por insuficiência de gas. Este resultado prático destaca que, para a viabilidade de um sistema similar em produção, mecanismos robustos de monitoramento de saldo e reabastecimento de fundos seriam componentes essenciais da infraestrutura operacional.

Em suma, a análise dos resultados demonstra que a arquitetura proposta é funcional, mas sua performance quando submetida a estresse extremo é limitada pela capacidade de processamento do nó Ganache. A estratégia de sincronização, embora essencial para a integridade do nonce, reforça o processamento ineficiente em cenários de paralelismo, comprometendo a escalabilidade da arquitetura.

Conclusão

Este trabalho demonstrou a viabilidade e os desafios da utilização de Smart Contracts na rede Ethereum para melhorar a segurança no registro de dispositivos IoT. A arquitetura proposta, que integra uma API RESTful com a blockchain, provou ser funcional e eficaz na garantia da integridade e auditabilidade do processo de autenticação de dispositivos. A modularização dos Smart Contracts contribuiu para a robustez e solidez da arquitetura proposta.

Os testes de estresse realizados com Apache JMeter revelaram aspectos importantes sobre o desempenho do sistema. Embora a solução garanta a correteza das transações e a prevenção

de problemas de nonce através da serialização, essa abordagem introduz um gargalo significativo na latência. A capacidade de processamento do nó Ganache, utilizado para simular a rede blockchain, foi identificada como o principal fator limitante para a escalabilidade em cenários de alta concorrência. Isso ressalta um trade-off inerente entre a garantia da integridade transacional e a velocidade de processamento em ambientes blockchain descentralizados.

Em suma, a pesquisa identifica o potencial da blockchain para fortalecer a segurança em ambientes de rede IoT, mas também aponta para a necessidade de otimização futuras. Para que soluções baseadas em blockchain sejam amplamente adotadas em ambientes IoT de larga escala, é importante desenvolver estratégias que diminuam os gargalos de desempenho, como a exploração de soluções que otimizem a capacidade de processamento das requisições (throughput) como também a otimização dos custos associados à criação de transações na rede da Ethereum.

Para superar as limitações encontradas neste estudo, como sugestão para trabalhos futuros é necessário resolver problemas de escalabilidade e performance por meio da utilização de outro nó blockchain mais robusto para que seja possível aumentar a taxa de processamento das requisições, e também testar a validade de utilizar um único contrato para realizar o registro dos dispositivos, a fim de verificar se ocorrerá uma mudança na latência ou throughput. Para solucionar o problema relacionado aos custos uma solução possível é a utilização das soluções Camada 2 (Layer-2) da rede Ethereum.

Para superar as limitações de performance identificadas sugere-se a substituição do nó de desenvolvimento Ganache por clientes de produção, para estabelecer uma linha de base de desempenho realista. Em conjunto, uma análise comparativa entre a atual arquitetura de contratos modulares e uma abordagem de contrato único (monolítico) para identificar o impacto do design do contrato na latência e na vazão do sistema. Adicionalmente, para mitigar os desafios de escalabilidade e custo de transação (gas), propõe-se a exploração e implementação da arquitetura em soluções de Camada 2 (Layer-2) da rede Ethereum, que são projetadas especificamente para oferecer alta vazão a custos reduzidos.

Referências

CARRION, Patricia; QUARESMA, Manuela. Internet da Coisas (IoT): Definições e aplicabilidade aos usuários finais. Human Factors in Design, Florianópolis, v. 8, n. 15, p. 049–066, 2019. DOI: 10.5965/2316796308152019049. Disponível em: <https://www.revistas.udesc.br/index.php/hfd/article/view/2316796308152019049>. Acesso em: 26 ago. 2025.

LEITE, Leandro Rogério Corrêa. Internet das Coisas (IoT): vulnerabilidades de segurança e desafios, 2019. Trabalho de conclusão de curso (Curso Superior de Tecnologia em Segurança da Informação) - Faculdade de Tecnologia de Americana, Americana, 2019.

RIBEIRO, Lucas; MENDIZABAL, Odorico. Introdução à Blockchain e Contratos Inteligentes. Apostila para iniciantes. Relatório Técnico do INE, Universidade Federal de Santa Catarina, 2021. Disponível em: <https://repositorio.ufsc.br/handle/123456789/221495>

BUTERIN, Vitalik. Ethereum: a next generation smart contract & decentralized application platform. [S. l.], 2013. Disponível em: <https://ethereum.org/en/whitepaper/>

Anexos

Figura 1: Diagrama da arquitetura do sistema.

Figura 2: Tela de requisição do Insomnia

Figura 3: Log da API de cadastro do sensor de Umidade

Figura 4: Média e desvio padrão do throughput médio das requisições

Figura 5: Média e desvio padrão do tempo de resposta dos endpoints

Tabela 1: Resultados do teste de carga

CÓDIGO: ET0512

AUTOR: CARLOS EDUARDO VALLE ROSA FILHO

ORIENTADOR: CHARLES ANDRYE GALVAO MADEIRA

TÍTULO: Experimentos com aprendizado por reforço profundo para construção de estratégias em ambientes complexos de jogos

Resumo

O projeto de pesquisa Aprendizado por Reforço em Jogos de Estratégia em Tempo Real tem como objetivo realizar experimentos de aprendizado por reforço em ambientes complexos de jogos, como Starcraft II. Isto é feito com auxílio da biblioteca URNAI — desenvolvida em paralelo aos experimentos — a qual procura facilitar a pesquisa no ramo de Aprendizado Profundo por Reforço em ambientes relacionados a jogos, oferecendo framework para experimentação de algoritmos. Neste documento, apresentarei o trabalho feito na bolsa entre 2024.2 e 2025.1, com foco em partes que tive maior participação.

Palavras-chave: Starcraft; Aprendizado por Reforço; IA; Aprendizado de Máquina; URNAI;

TITLE: Experiments with deep reinforcement learning for building strategies in complex game environments

Abstract

The Reinforcement Learning in Real-Time Strategy Games research project aims to conduct reinforcement learning experiments in complex game environments, such as Starcraft II. This is done using the URNAI library—developed in parallel with the experiments—which aims to facilitate research in Deep Reinforcement Learning in game-related environments by providing a framework for algorithm experimentation. In this document, I will present the work done during the fellowship between 2024.2 and 2025.1, focusing on the parts in which I had the greatest involvement.

Keywords: Starcraft; Reinforcement Learning; AI; Machine Learning; URNAI;

Introdução

O Aprendizado por Reforço (RL) é uma área de conhecimento interessada no comportamento de um agente inteligente em um ambiente dinâmico de maneira a maximizar uma recompensa. Aplicações incluem, mas não são limitadas a: Veículos Autônomos, Otimização de Sistemas em saúde, finanças e indústria, Algoritmos de Recomendação e Automação de Robôs [7, 13].

Esta área já mostrou excelentes resultados em cenários de jogos como Xadrez e Go, derrotando jogadores profissionais em ambos [12]. Entretanto, aplicações de RL em situações mais próximas à realidade — a locomoção de um robô humanoide, por exemplo [1] — apresentam um desafio significativamente maior devido à alta complexidade do cenário

modelado. Visto que, diferentemente dos jogos citados, a quantidade de ações disponíveis é maior, bem como a quantidade de informações no estado do agente. Dessa forma, a pesquisa nessa área necessita de cenários cada vez mais complexos para progredir.

O trabalho feito na bolsa durante o período de 2024.2 e 2025.1 consistiu principalmente em experimentos de RL no cenário de Starcraft II, cuja complexidade atende às necessidades do trabalho de pesquisa. Starcraft II é um jogo de estratégia em tempo real — onde os jogadores devem orquestrar um exército e meios de produção de forma a eliminar o oponente — o qual ganhou popularidade nos últimos anos no contexto de aprendizado por reforço [14]. Os experimentos com este cenário foram feitos por meio da biblioteca URNAI [9], uma biblioteca modular de aprendizado por reforço profundo.

Além disso, parte do trabalho focou no desenvolvimento da própria URNAI. Esta está em sua segunda versão, pois a anterior apresentou dificuldades em relação aos experimentos, mostrando inconsistências nos dados gerados, o que dificultou uma análise aprofundada dos resultados obtidos. Dessa forma, decidiu-se desenvolver uma nova versão, com um maior formalismo e uma estratégia de testes de código, garantindo uma maior robustez.

Metodologia

O código foi feito na linguagem Python em sua versão 3.10, com o versionamento feito pela ferramenta Git [4] e hospedagem na plataforma GitHub [5]. Além da URNAI, foram utilizadas outras duas principais bibliotecas: PySC2 [8] para o ambiente em Starcraft II e Stable Baselines 3 [2] a qual fornece implementações confiáveis de algoritmos de aprendizado. Ademais, a execução foi feita em uma imagem da plataforma Docker [3] para a distribuição Linux.

Quanto à gestão dos experimentos, esta foi feita por meio da plataforma Weights & Biases [15], a qual proporciona uma interface gráfica que mantém informações sobre todas as execuções registradas. Nela, também é possível uma visualização customizada dos gráficos gerados a partir dos dados dos experimentos, permitindo, por exemplo, a escolha da função de suavização do gráfico e a alteração da escala e cor.

Os experimentos abordaram três mini-games — cenários simplificados — propostos pela PySC2: CollectMineralShards, no qual o agente deve aprender a coletar o máximo de minerais espalhados pelo mapa em um limite de tempo; DefeatRoaches, no qual o agente deve aprender a eliminar ondas de inimigos; BuildMarines, no qual o agente deve produzir o maior quantidade possível da unidade “marine” em um limite de tempo. Nestes, foram testadas principalmente variações no espaço de ação, espaço de observação e configuração de recompensa do agente — três fatores essenciais no aprendizado por reforço — em busca de uma melhor performance. O algoritmo de aprendizado utilizado foi o PPO, uma escolha popular para problemas de RL [11].

Resultados e Discussões

O primeiro mini-game testado foi o CollectMineralShards. Este, apesar de não apresentar um resultado (Figura 1) em par com aquele publicado pelos criadores do cenário [14], mostrou que a segunda versão da biblioteca URNAI estava funcional, visto que observou-se por meio dos gráficos e testes realizados, uma evolução da performance do agente.

Com relação aos outros dois mini-games, inicialmente, apresentaram inconsistências entre o desempenho do agente no treinamento em relação aos testes — com um bom desempenho no primeiro e péssimo no segundo — algo não esperado. Após um período de depuração, o qual envolveu a execução de outros cenários mais simples, como aqueles disponíveis na interface Gymnasium [6], e análises para o comportamento do agente para situações específicas, conclui-se que a irregularidade observada se devia a dois fatores: falta de restrições de ações inválidas e atribuição incorreta da recompensa a uma certa ação. Ambos fatores prejudicam o aprendizado do agente, o que poderia explicar a disparidade de comportamento exibida nos testes. Tais problemas se expressaram mais fortemente no cenário BuildMarines, logo optou-se por torná-lo o foco das tentativas de solução.

A respeito das restrições de ações inválidas, adotou-se como solução a utilização da versão MaskablePPO do algoritmo PPO, a qual permite restringir quais ações são disponíveis em determinado momento para o agente. Esta versão é disponibilizada pela biblioteca RL Baselines3 Zoo [10], derivada da Stable Baselines 3. Tal medida melhorou o desempenho do agente nos testes, mas não eliminou a disparidade em relação ao treinamento (Figuras 2 e 3). Vale notar que também foi testada a penalização de ações inválidas como uma alternativa, mas esta não foi tão efetiva quanto.

Já para a atribuição incorreta da recompensa a uma certa ação, houveram duas soluções propostas: fornecer um histórico de ações ao agente em seu espaço de observação e um sistema de recompensas retardadas. A primeira solução melhorou o desempenho do agente, mas não resolveu o problema por completo. Quanto à segunda solução, nota-se que esta se encaixa bem em situações onde as ações não são instantâneas — ou seja, as suas recompensas também não são — que é o caso do cenário BuildMarines. Entretanto, descobriu-se que esta abordagem não é compatível com a Stable Baselines 3 — responsável pela parte relevante a esse problema no código — a qual, por design, sempre associa uma recompensa à ação executada naquele instante. Por fim, tentou-se outra abordagem: a atribuição da recompensa de um episódio inteiro apenas em seu final. Isto, em teoria, garante o aprendizado do agente, porém mais lentamente. Atualmente, os experimentos com esta configuração ainda não foram finalizados, mas por hora os resultados são em sua maior parte positivos.

Conclusão

Os experimentos com a biblioteca URNAI no cenário de Starcraft II, apesar de ainda não produzirem resultados desejados em sua maior parte, demonstram o funcionamento da biblioteca URNAI, cujo desenvolvimento também se caracteriza como parte do trabalho. Além de garantir uma base para experimentos futuros, com a experiência adquirida na resolução dos desafios encontrados.

Para trabalhos futuros, tem-se em mente a experimentação em mais cenários relacionados à Starcraft II e outros jogos com cenários complexos. Outro objetivo é aprimorar o desempenho nos cenários já testados.

Referências

1. Deep Reinforcement Learning for Robotic Bipedal Locomotion: A Brief Survey. Disponível em: <<https://arxiv.org/html/2404.17070v2>>. Acesso em: 18 ago. 2025.
2. DLR-RM/stable-baselines3. Disponível em: <<https://github.com/DLR-RM/stable-baselines3>>.
3. DOCKER. Enterprise Application Container Platform | Docker. Disponível em: <<https://www.docker.com/>>.
4. Git. Disponível em: <<https://git-scm.com>>.
5. GitHub. Disponível em: <<https://github.com>>.
6. Gymnasium Documentation. Disponível em: <<https://gymnasium.farama.org/>>.
7. LI, Y. REINFORCEMENT LEARNING APPLICATIONS. 2019. [s.l: Site (Medium).]. Disponível em: <<https://arxiv.org/pdf/1908.06973>>. Acesso em: 24 ago. 2024.
8. PySC2. Disponível em: <<https://github.com/google-deepmind/pysc2>>.
9. URNAI. Disponível em: <<https://github.com/UFRN-URNAI>>. Acesso em: 24 ago. 2024.
10. RL Baselines3 Zoo — Stable Baselines3 2.3.0a1 documentation. Disponível em: <https://stable-baselines3.readthedocs.io/en/master/guide/rl_zoo.html>.
11. SCHULMAN, J. et al. Proximal Policy Optimization Algorithms. Disponível em: <<https://arxiv.org/abs/1707.06347>>.
12. SILVER, D. et al. Mastering the game of Go with deep neural networks and tree search. Nature, v. 529, n. 7587, p. 484–489, jan. 2016.
13. SIVAMAYIL, K. et al. A Systematic Study on Reinforcement Learning Based Applications. Energies, v. 16, n. 3, p. 1512, 1 jan. 2023. [s.l: Periódico (Energies)].
14. VINYALS, O. et al. StarCraft II: A New Challenge for Reinforcement Learning. arXiv:1708.04782 [cs], 16 ago. 2017.
15. Weights & Biases – Developer tools for ML. Disponível em: <<https://wandb.ai/site>>.

Anexos

Figura 2

Figura 3

Figura 1

CÓDIGO: ET0788

AUTOR: JADY LIMA DA SILVA

COAUTOR: GABRIEL ROCHA DOS SANTOS

COAUTOR: PATRICK CESAR ALVES TERREMATTE

COAUTOR: HOSANA IASMIN CASTRO DOS SANTOS

ORIENTADOR: RICARDO JOSE MATOS DE CARVALHO

TÍTULO: Análise da usabilidade e da satisfação do usuário do sistema web NOAH.

Resumo

Este estudo teve como objetivo inicial analisar e reconstruir o sistema web Noah de apoio ao gerenciamento de riscos de desastres da cidade de Natal-RN, desenvolvido em 2019. Durante as análises, verificou-se a necessidade de modernizar a estrutura deste sistema, desenvolvendo um novo sistema web, por meio de novas tecnologias, utilizando, especialmente, a Inteligência Artificial, para auxiliar as atividades de gerenciamento do órgão municipal de Proteção e Defesa Civil, que consiste, dentre outras, em coordenar ações, receber e emitir informações e se comunicar com a população e outros agentes públicos, antes, durante e depois da ocorrência de desastres. Além disso, foi conduzido um estudo de caráter exploratório e prático, a fim de identificar, quantificar e mapear as ocorrências (ameaças ou perigos, desastres, etc.) registradas nos últimos 13 anos (2013-2025) pelo referido órgão, trazendo uma análise crítica acerca dos dados observados. Os resultados indicaram a concentração de ocorrências em determinadas zonas da cidade, além de uma relação direta com os períodos chuvosos, sugerindo a necessidade de mais ações preventivas em infraestrutura. Conclui-se que estratégias de gerenciamento de riscos e programas de capacitação da população são essenciais para melhorar a resposta em situações de emergências. Também, que as tecnologias desenvolvidas produzem dados que auxiliam as análises, tomadas de decisão e ações dos agentes do órgão para mitigar riscos e salvar vidas.

Palavras-chave: Desastres. Prevenção. Gerenciamento de Riscos. Ocorrências.

TITLE: Usability and User Satisfaction Analysis of the NOAH Web System

Abstract

This study initially aimed to analyze and reconstruct the Noah web system, designed in 2019 to support disaster risk management in the city of Natal-RN. During the analyses, the need to modernize the system's structure was identified, leading to the development of a new web system based on emerging technologies, with a particular emphasis on Artificial Intelligence, to support the activities of the Municipal Protection and Civil Defense Agency. These activities include, among others, coordinating actions, receiving and disseminating information, and communicating with the population and other public entities before, during, and after disaster events. In addition, an exploratory and practical study was conducted to identify, quantify, and map the occurrences (threats, hazards, disasters, etc.) recorded by the agency over the last 13 years (2013–2025), providing a critical analysis of the observed data. The results indicated a concentration of occurrences in specific areas of the city, as well as a

direct relationship with rainy periods, suggesting the need for further preventive infrastructure measures. It is concluded that risk management strategies and community training programs are essential to improve emergency response. Furthermore, the developed technologies produce data that support analyses, decision-making processes, and the actions of the agency's personnel in mitigating risks and saving lives.

Keywords: Disasters. Prevention. Risk Management. Occurrences.

Introdução

Este trabalho tem como objetivo apresentar e discutir os registros do órgão municipal de proteção e defesa civil (OMPDC) de Natal-RN referentes às ameaças e danos ocorridos entre 2013 2025 por conta de incidentes e desastres no município, e descrever o desenvolvimento do sistema web Noah, para propiciar a comunicação entre este órgão e a população e a coordenação das ações de gerenciamento de riscos de desastres (GRD).

O crescimento acelerado da população mundial resultou em uma maior concentração de pessoas em áreas urbanas, frequentemente de forma desorganizada e sem planejamento, gerando novos desafios no GRD.

Segundo o United Nations Office for Disaster Risk Reduction (UNDRR), desastre é “uma interrupção grave do funcionamento de uma comunidade ou sociedade em qualquer escala devido a eventos perigosos que interagem com condições de exposição, vulnerabilidade e capacidade, levando a um ou mais dos seguintes: perdas e impactos humanos, materiais, econômicos e ambientais” (UNDRR, 2017). Para Quarantelli in Perry & Quarantelli (2005, p. 343-347), o desastre é um fenômeno social resultante das ações dos seres humanos e de suas sociedades; se não há consequência social, não há desastre. Estima-se que a média de gastos relacionados a perdas causadas por estes eventos chegue a US\$ 260 bilhões em 2015 para US\$ 414 bilhões em 2030, segundo projeções do UNISDR (2019).

Considerando esse contexto, o Sistema Web Noah foi desenvolvido e registrado no ano de 2019 no Instituto Nacional da Propriedade Industrial buscando auxiliar a defesa civil no GRD em Natal, visando, principalmente, à diminuição de perdas materiais e impactos humanos. O sistema, em sua versão atual, está integrado a um projeto de pesquisa mais amplo, com outros núcleos de pesquisas, voltado, especialmente, para as aplicações técnicas de Inteligência Artificial (IA) no contexto da gestão de riscos de desastres. Contudo, a presente pesquisa tem como recorte principal o sistema Web Noah, que constitui a ramificação do projeto ligada ao apoio das atividades do OMPDC da cidade de Natal.

Este estudo fez-se necessário para aprimorar o sistema web Noah que estava descontinuado desde 2021 em virtude do falecimento de um dos coordenadores do projeto em decorrência da pandemia da COVID-19. Além da reativação, buscou-se incorporar tecnologias atuais que possibilitam melhorias em sua estrutura, alinhando-o às tendências científicas e ampliando sua eficiência no apoio ao GRD Assim, embora inserido em um projeto de escopo maior, este trabalho concentra-se na análise e evolução do sistema Web Noah.

Diante deste contexto, formulou-se a questão norteadora de pesquisa: o desenvolvimento e utilização de um sistema web, através de tecnologias emergentes, como a Inteligência Artificial, pode auxiliar os agentes do OMPDC no GRD, tais como, a tomar melhores decisões, a se comunicar entre si e com a população, e a agir melhor com base em dados produzidos por esta tecnologia?

O objetivo geral do trabalho consiste em analisar e reconstruir o sistema Web Noah, modernizando sua estrutura através de novas tecnologias, de modo a torná-lo mais eficiente no apoio às ações de gestão de riscos e desastres em Natal. Como objetivos específicos, destacam-se: (i) reconstruir as funcionalidades essenciais do sistema; (ii) identificar e mapear as categorias de ocorrências predominantes de perigos ou ameaças, desastres e vistorias, interdições e desinterdições realizadas pelo OMPDC entre os anos de 2013 e 2025 em Natal; (iii) avaliar as ocorrências mais frequentes em determinados períodos e zonas; e (iv) alinhar o desenvolvimento e a visualização dos dados de forma a subsidiar as tomadas de decisão e ações do OMPDC.

Pretende-se, com isto, contribuir com o GRD e com a área de tecnologia de informação e comunicação aplicada ao GRD. Espera-se disponibilizar esta ferramenta ao OMPDC, contribuir para a segurança das populações vulneráveis a desastres e incentivar novas ações no contexto do GRD.

Metodologia

A presente pesquisa caracteriza-se como aplicada, por buscar a solução prática de problemas específicos, e de caráter exploratório, por possibilitar maior familiaridade com o objeto estudado, conforme Marconi e Lakatos (2003). O estudo concentra-se na reconstrução do sistema web Noah, por meio da utilização e do mapeamento de novas tecnologias.

Inicialmente, realizou-se o contato com os desenvolvedores anteriores do sistema web Noah para compreender a implementação original do sistema. Foram analisadas o funcionamento e a organização da aplicação. Por meio dessa etapa, foi recuperada uma versão antiga da aplicação desenvolvida exclusivamente em PHP no backend com uso do Bootstrap para recursos visuais. No entanto, não se teve acesso a dados acerca de como foi estruturado o sistema de banco de dados, o que representou uma limitação considerável para a continuidade do trabalho.

Em sequência, com uma versão de código em mãos, buscou-se compreender como seria possível substituir o banco de dados original atendendo às necessidades funcionais do sistema web Noah. Durante esse processo, foi observada a possibilidade de migração do sistema original Noah, de modo a favorecer a reestruturação e manutenção futura da aplicação.

Dessa forma, optou-se pela adoção do framework Laravel, devido a sua escalabilidade, robustez e facilidade de integração a banco de dados, alinhando-se às boas práticas de engenharia de software, que indicam frameworks modernos como essenciais para a manutenção e evolução contínua do sistema (SOMMERVILLE, 2016). Constatou-se que o framework escolhido não apenas apresenta vantagens em termos de manutenção, mas, em documentação abrangente e a atividade constante da comunidade de desenvolvedores na evolução contínua do sistema (LARAVEL, 2025).

Com a definição do framework, foi feito o uso da estrutura original recuperada do sistema web Noah, por meio das migrations, foi gerado o esquema de atendimento das necessidades essenciais do sistema. As etapas seguintes da pesquisa foram voltadas para o estudo das tecnologias escolhidas e o desenvolvimento de uma nova aplicação Noah, semelhante ao sistema anterior, com foco nas melhorias proporcionadas pelo framework, como o sistema de

autenticação e integração ao banco de dados e o uso do Bootstrap atualizado para melhores apresentações.

Com a versão estável do sistema, iniciou-se a etapa seguinte, voltada à obtenção e análise de dados. O contato com gestores e agentes do órgão de proteção e defesa civil municipal foi fundamental nesse momento, pois era necessário compreender melhor as necessidades do cliente e definir o Mínimo Produto Viável (MVP). Essa definição seguiu a abordagem de desenvolvimento ágil de software, conforme proposta por Eric Ries (2011), que descreve o MVP como “a versão do produto que permite uma volta completa do loop Build-Measure-Learn com o mínimo esforço e o menor tempo de desenvolvimento possível”.

Por meio da base de dados cedida pela defesa civil, atuou-se principalmente com o foco de inserir os dados no sistema. No primeiro momento desta etapa, foi realizada a limpeza dos dados originais e a geração de dados de coordenadas genéricas, seguindo as práticas de anonimização dos dados, considerando a segurança dos dados reais, juntamente com a separação dos dados que necessitavam ser apresentados e, também, avaliou-se, definiu-se e planejou-se a melhor forma de exibi-los. Com as informações separadas e organizadas, foi possível desenvolver uma aplicação na linguagem de programação, Python, Noah Dados, que atuou enviando os dados ao endpoint da aplicação responsável por salvar no sistema todas as ocorrências obtidas, do ano de 2013 até março de 2025.

Também, foi criada uma estrutura além do Noah Dados, o Noah Chatbot, um trabalho de extensão integrado ao Projeto Noah, com o objetivo de aprimorar a comunicação entre a comunidade e a defesa civil. A proposta central consistia em que o sistema web fosse capaz de atender a especificações que permitissem sua integração ao chatbot, disponibilizado no WhatsApp, possibilitando a interação direta com os usuários em situação de risco e o registro automático das ocorrências na aplicação web.

Com os dados das ocorrências anteriores, fornecidos pelo órgão municipal de proteção e defesa civil, foram definidas estratégias para o mapeamento e a visualização das informações no sistema. Inicialmente, com foco principal nos dados, considerando zonas ou regiões geográficas (norte, sul, leste, oeste), categorias de ocorrência, bem como a descrição e o mapeamento dos locais com maior volume de registros ao longo dos períodos analisados.

Para a implementação da visualização dos dados no dashboard do sistema, optou-se pelo uso da biblioteca ApexCharts, integrada ao JavaScript, devido a sua flexibilidade para a criação de gráficos interativos e a sua facilidade de integração a sistemas web. Essa escolha possibilitou a geração de gráficos dinâmicos para a análise de categorias e variações ao longo do tempo. De forma complementar, utilizou-se a biblioteca Leaflet para a construção do mapa de calor, permitindo a identificação das áreas mais críticas quanto ao registro de ocorrências. A adoção dessas bibliotecas segue os princípios de visualização de dados, conforme Tufte (2001) e Evergreen (2019), priorizando clareza, compreensão rápida das informações e suporte à tomada de decisão pelos órgãos competentes.

Dessa forma, a metodologia contemplou a análise do sistema original, a reestruturação do banco de dados e migração para novas tecnologias, assim como a integração do sistema aos novos módulos e implementação de visualizações de dados interativas, garantindo eficiência para com a gestão de riscos de desastres na cidade. Todas as etapas foram conduzidas buscando o cumprimento dos objetivos de atualização e aprimoramento do sistema Web Noah, atendendo às necessidades do órgão municipal de Proteção e Defesa Civil e possibilitando análises fundamentadas nas ocorrências anteriores.

Resultados e Discussões

Nesta seção, apresentaremos os resultados obtidos com o desenvolvimento do sistema Web Noah atrelado aos objetivos que modelaram a presente pesquisa. Os dados, fornecidos pelo órgão de Proteção e Defesa Civil de Natal, apresentados na Figura 1, abrangem o período de 2013 a 2025 e oferecem informações relevantes para compreender a distribuição das ocorrências de vistorias, interdições e desinterdições de setores de riscos na cidade de Natal - RN.

Iniciamos as análises, considerando entender como estão distribuídas estas ocorrências considerando séries temporais. O Figura 1 nos mostra dados relacionados às ocorrências ao longo dos meses de 2024. Quando comparados aos anos anteriores, notou-se um padrão, entre o segundo e terceiro trimestre destes anos, os números de ocorrências registradas apresentaram aumento acentuado.

Observou-se que essa tendência está fortemente associada aos períodos chuvosos da cidade (meses de maio a agosto), evidenciando a relação direta entre situações climáticas e registros da defesa civil.

Os estudos seguintes concentraram-se na observação acerca das zonas mais afetadas pelas ocorrências e quais tipos de eventos estão associados a elas. Verificou-se uma repetição de comportamento, semelhante à análise anterior, na qual, em sua maioria, os registros estavam relacionados à Zona Norte da cidade, onde concentram-se as lagoas de captação de água pluvial que vêm apresentando transbordamentos, inundando vias, domicílios e empresas. Além disso, foi possível notar não apenas a maior concentração de ocorrências, mas, também, que os casos de interdição estavam mais relacionados ao local, como mostrado no Figura 2.

Os dados evidenciam uma maior vulnerabilidade da Zona Norte com 710 registros dos 953 ocorridos no mesmo ano, o que vem se repetindo nos últimos 5 anos, possivelmente associada não apenas às condições climáticas, mas também a fatores de infraestrutura e socioeconômicos. Esse cenário pode sugerir a necessidade de maiores investimentos e ações direcionadas à região, de modo a reduzir as ocorrências.

Seguindo com as análises dos gráficos obtidos, foi gerada uma visualização (Anexo 3) que permitisse entender quais os 10 principais laudos, ordenados de forma decrescente, registrados durante cada ano.

Observa-se na Figura 3 que os tipos mais recorrentes estão associados às Vistorias. Esse resultado reforça a interpretação inicial obtida pela análise da Figura 1, em que os tipos podem estar relacionados a problemas gerados nos períodos chuvosos da região, como o transbordamento de lagoas com cerca de 538 registros, além de problemas relacionados a crateras em vias e alagamentos de imóveis, que, somados, ficam em torno de 61 ocorrências.

Visando entender melhor as dimensões considerando todos os tipos de ocorrências do ano, elaborou-se um TreeMap (Figura 4), desta forma, foi possível entender claramente as proporções e representatividade de cada uma delas. No primeiro trimestre de 2025, observa-se novamente a predominância de vistorias, seguidas por interdições, padrão consistente com os resultados de 2024 (Figura 3).

Entre as vistorias, destacam-se, especialmente, os problemas de infraestrutura urbana, como crateras em vias e alagamentos em imóveis, o que reforça a recorrência de vulnerabilidades já apontadas em períodos anteriores.

Considerando melhor compreender a distribuição das ocorrências por determinadas zonas no ano, elaborou-se um Diagrama de Sankey como mostrado na Figura 5, visando a melhor visualização dos fluxos. A ideia central é que, por meio desse rastreamento, a defesa civil pode não apenas compreender quais locais foram mais afetados, mas, também, propor ações a fim de amenizar os impactos ocasionados, além de realizar atividades preventivas.

A análise do primeiro trimestre de 2025 evidencia um fluxo maior de registros na Zona Norte, tendência que reforça o padrão já observado em 2024 (Figura 2). Esse resultado sugere a persistência de vulnerabilidades atreladas à região, que demandam medidas para redução das ocorrências recorrentes.

O mapa de calor foi fundamental para visualizar a intensidade dos registros em diferentes regiões. A Figura 6 apresenta a distribuição das ocorrências na cidade de Natal ao longo de 13 anos de dados obtidos, destacando os pontos de maior concentração. Essa visualização fornece subsídios importantes para a Defesa Civil, permitindo identificar áreas prioritárias e planejar ações de forma mais direcionada.

De modo geral, os resultados obtidos apresentam padrões consistentes acerca das ocorrências, especialmente relacionadas às condições climáticas e às características específicas socioeconômicas de cada região. Os dados trazem evidências que fornecem uma base consistente para discussões futuras e para o desenvolvimento de estratégias mais eficazes de monitoramento e intervenção.

Conclusão

Os resultados demonstram a necessidade de ações para fortalecer o gerenciamento de riscos de desastres, abrangendo desde o mapeamento das ocorrências até a capacitação da população sobre como agir em situações de emergência.

O sistema web Noah desenvolvido respondeu à questão de pesquisa formulada, atingiu os objetivos formulados e apresentou boa aceitação por parte dos gestores e agentes do órgão municipal de proteção e defesa civil. Espera-se, como trabalho futuro, expandi-lo para outros órgãos estratégicos, como outros órgãos municipais e estaduais de proteção e defesa civil e a Secretaria de Saúde Pública (Sesap), de modo a ampliar sua aplicabilidade.

Apesar das contribuições obtidas, o presente estudo apresenta limitações acerca da coleta de dados em campo, o que restringe a validação dele em situações reais. Sabendo disto, sugerem-se novas pesquisas que busquem realizar testes aplicados a grupos predefinidos (membros de populações vulneráveis a desastres), de forma a avaliar o desempenho do sistema web Noah em cenários práticos, como, por exemplo, em exercícios simulados com cenário postulado de evacuação de emergência de uma população de um setor de uma localidade afetada por um desastre.

A partir desses experimentos, será possível apresentar melhorias tanto na arquitetura do sistema quanto nos recursos voltados à visualização de dados, como a inclusão de filtros

temporais, além de uma reestruturação do banco relacional, assegurando maior eficiência e confiabilidade do sistema web Noah.

Referências

APEXCHARTS. **ApexCharts – Interactive JavaScript Charts for your Website**. Disponível em: <https://apexcharts.com>. Acesso em 21 ago. 2025.

LARAVEL. **Laravel: The PHP Framework for Web Artisans**. Disponível em: <https://laravel.com>. Acesso em: 20 ago. 2025.

MARCONI, Marina de Andrade; LAKATOS, Eva Maria. **Fundamentos de metodologia científica**. 6. ed. São Paulo: Atlas, 2003.

PERRY, R. W; QUARANTELLI, E. L.. **What is a Disaster? New answers to old questions**. USA: IRCD, 2005.

RIES, Eric. **The Lean Startup: How Today's Entrepreneurs Use Continuous Innovation to Create Radically Successful Businesses**. New York: Crown Business, 2011.

SOMMERVILLE, Ian. **Engenharia de Software**. 10. ed. Pearson, 2016.

TUFTE, Edward R. **The Visual Display of Quantitative Information**. 2. ed. Cheshire: Graphics Press, 2001.

United Nations Office for Disaster Risk Reduction (UNDRR). 2017. **The Sendai Framework Terminology on Disaster Risk Reduction."Disaster"**. Disponível em: <https://www.undrr.org/terminology/disaster>. Acesso em: 20 ago. 2025.

United Nations Office for Disaster Risk Reduction (UNDRR). **Sendai Framework for Disaster Risk Reduction 2015-2030**. Geneva: UNDRR, 2015. Disponível em: <https://www.undrr.org/publication/sendai-framework-disaster-risk-reduction-2015-2030>. Acesso em: 20 ago. 2025.

Anexos

Figura 1 - Distribuição de Ocorrências por Meses do Ano

Figura 2 - Distribuição de Ocorrências por Zonas

Figura 3 - Tipos de Ocorrências mais Frequentes em 2024

Figura 4 - Distribuição Geral das Ocorrências por Ano

Figura 5 - Fluxo de Ocorrências por Região

Figura 6 - Mapa de Distribuição Geral de Ocorrências em Natal

CÓDIGO: ET1060

AUTOR: JULIANA FREIRE PEQUENO DE SANTIAGO CARVALHO

ORIENTADOR: RENAN CIPRIANO MOIOLI

TÍTULO: Uso do eyetracking como metodologia para compreensão do Processamento de Expressões Idiomáticas pelo cérebro humano

Resumo

Este estudo investiga o uso do rastreamento ocular (eyetracking) para entender melhor o processamento de expressões idiomáticas, contribuindo para o desenvolvimento de modelos de Processamento de Linguagem Natural (PLN) mais eficientes. Explora-se como os movimentos oculares, compostos por sacadas e fixações, oferecem insights sobre os processos cognitivos durante a leitura, conforme destacado por Rayner (2009). O projeto está associado à iniciativa MIA, liderada pela professora Aline Villavicencio, e visa integrar pistas linguísticas e cognitivas para melhorar a compreensão de expressões idiomáticas por modelos computacionais. O experimento construído neste trabalho foi o de reconstruir o programa para coletar dados oculares durante a leitura de expressões idiomáticas em inglês britânico. Os dados coletados contribuirão à realização de estudos comparando o entendimento de falantes nativos e não nativos de inglês, fornecendo insights valiosos para a pesquisa em psicolinguística e neurociência. Além disso, o uso de softwares como Psychopy facilita o acesso ao experimento e à sua construção, já que é open source. Enquanto a colaboração com o grupo de Sheffield potencializa os resultados em um contexto global.

Palavras-chave: Eyetracking; Psychopy; Tobii; Expressões Idiomáticas; Modelos PLN;

TITLE: Use of eye tracking as a methodology for understanding the processing of idiomatic expressions by the human brain

Abstract

This study investigates the use of eye tracking to better understand the processing of idiomatic expressions, contributing to the development of more efficient Natural Language Processing (NLP) models. It explores how eye movements, consisting of saccades and fixations, offer insights into cognitive processes during reading, as highlighted by Rayner (2009). The project is associated with the MIA initiative, led by Professor Aline Villavicencio, and aims to integrate linguistic and cognitive cues to improve the understanding of idiomatic expressions by computational models. The experiment constructed in this work was to rebuild the program to collect eye data while reading idiomatic expressions in British English. The data collected will contribute to studies comparing the understanding of native and non-native English speakers, providing valuable insights for research in psycholinguistics and neuroscience. In addition, the use of software

such as Psychopy facilitates access to the experiment and its construction, as it is open source. Meanwhile, collaboration with the Sheffield group enhances the results in a global context.

Keywords: Eye tracking; Psychopy; Tobii; Idiomatic expressions; PLN models;

Introdução

A psicolinguística é uma disciplina da psicologia cognitiva que investiga a relação entre linguagem e mente, analisando como os humanos processam informações e compreendem a linguagem. A técnica conhecida como Eyetracking, que mede os movimentos oculares durante a leitura, é amplamente utilizada neste campo, pois permite explorar diversos aspectos do processo de percepção visual e do processamento textual pelo cérebro humano (Rayner et al., 1989; Hollenstein, Barrett e Björnsdóttir, 2022). O psicólogo cognitivo Keith Rayner, um dos pioneiros no uso do rastreamento ocular, revisou a aplicação desse método sob diferentes perspectivas, avaliando como os indivíduos reagem ao que leem em textos ou visualizam em imagens. Conforme suas conclusões, os dados de Eyetracking oferecem insights valiosos sobre os processos cognitivos que ocorrem continuamente durante a leitura (Rayner, 2009).

Os movimentos oculares são compostos por dois elementos principais: sacadas e fixações. Sacadas são movimentos curtos ou rotativos, enquanto fixações são os intervalos em que os olhos permanecem parados, permitindo o processamento de novas informações (Rayner, 2009, p. 1458). As sacadas podem ser classificadas de acordo com a direção da leitura, sendo a regressão um tipo de sacada que movimenta os olhos para trás no texto. Esta é a terceira componente mais relevante dos movimentos oculares na leitura, ocorrendo em cerca de 10-15% do tempo em leitores ávidos (Rayner, 2009, p. 1458). Quando aplicado à leitura, o rastreamento ocular sincroniza as sacadas e fixações com a posição do texto, possibilitando a avaliação das decisões de fixação visual em relação à localização e ao tempo gasto em cada ponto. A utilização de dados de rastreamento ocular é igualmente significativa para o desenvolvimento de modelos de Processamento de Linguagem Natural (PLN), visando aprimorar seus resultados. Os modelos de controle do movimento ocular durante a leitura ajudam a entender para onde os leitores direcionam seu olhar e quanto tempo permanecem em cada palavra (Rayner, 2009, p. 1487).

Expressões idiomáticas são frases compostas por um conjunto de palavras cujo significado não pode ser deduzido a partir do sentido literal de seus componentes. Em outras palavras, seu significado não é diretamente derivado dos significados dos elementos que as compõem (Expressão Idiomática, 2024). No contexto social, um aspecto fundamental da comunicação humana é o uso de expressões idiomáticas, que requerem uma compreensão contextual para garantir que a interação seja fluida e eficaz. No entanto, estudos recentes indicam que os modelos de PLN enfrentam desafios significativos ao lidar com expressões idiomáticas, o que compromete a naturalidade e precisão na interpretação linguística (Mi, Villavicencio & Moosavi, 2024).

De acordo com Rayner (2009), uma vantagem dos dados de movimento ocular é que eles oferecem uma boa indicação dos processos cognitivos a cada movimento durante a leitura, como já mencionado anteriormente. Rayner realizou um estudo que reuniu diversos trabalhos relacionados a movimentos oculares e atenção na leitura, percepção de cenas e busca visual.

Neste artigo, ele compilou várias abordagens utilizadas em conjunto com o equipamento Eyetracker, além de reunir descobertas comprovadas por meio desses estudos paralelos.

Devido à sua falta de composicionalidade semântica, as expressões idiomáticas são mentalmente representadas como longas palavras morfologicamente complexas e são reconhecidas através do mesmo processo de recuperação que ocorre durante o reconhecimento de palavras (Tabossi et al., 2009). Esse processo de recuperação começa assim que a primeira palavra da expressão idiomática é identificada, ocorrendo em paralelo ao processamento do sentido literal. No entanto, a computação do significado é um processo mais lento do que a recuperação.

Metodologia

Estudo - Modelos Profundos de Redes Neurais Aprimorados para Capturar Nuances da Linguagem como Expressões Idiomáticas

Primeiramente, foi realizada uma adaptação e validação da rotina do experimento original para o software que estará sendo utilizado neste projeto. A estrutura e ordem dos eventos dentro da rotina foram mantidos, assim como também todas as frases utilizadas e as perguntas de compreensão. É importante ressaltar que o experimento será realizado completamente em inglês.

A rotina funciona da seguinte maneira: Inicialmente o participante irá responder um formulário com perguntas básicas, contendo seu nome, seu nível de proficiência em inglês e sua idade. Após isso, o procedimento dará início.

1. O participante realizará a calibração do Eyetracker seguindo uma rotina de ajuste personalizada para o seu padrão visual. Utilizando um array de 9 pontos localizados na tela.
2. Uma rotina de teste será iniciada, contendo 6 expressões idiomáticas com mais de uma palavra em exemplos de frases.
 - A fonte utilizada é a Open Sans.
 - A letra tem o tamanho de 1.50% da largura da tela (`win.size[0]`)
 - A posição do estímulo está centrada no meio da tela (0,0).
 - O sistema de unidade utilizado em todo o experimento é o de pixels
3. A cada frase, o participante lerá, e, quando se sentir confortável, vai pressionar qualquer tecla do teclado para passar para a próxima
4. Ao finalizar as 6 palavras, uma pergunta de compreensão será apresentada na tela. Onde o participante terá que responder de [1,5] - caso seja uma pergunta escalar - e [0,4] - caso seja uma pergunta numeral.
5. Após a finalização da etapa de teste. O participante será avisado que o experimento será iniciado.
6. A cada 4 frases que sejam apresentadas, o participante terá de responder uma pergunta de compreensão, selecionada aleatoriamente.

7. Essa rotina se repetirá até que o usuário alcance metade do banco de dados, que possui cerca de 198 expressões idiomáticas. Ao alcançar a metade do banco de dados, o participante terá um tempo para repousar antes de prosseguir.

Devido a quantidade de expressões idiomáticas, o experimento pode tomar um tempo considerável de 1h para ser realizado. Dado isso, a pausa no meio do experimento pode contribuir para que o participante não sinta fadiga durante a sua experiência.

Os materiais utilizados foram:

- Tobii Pro Spark Eyetracker: É um eyetracker de 60hz da empresa Tobii, é um aparelho fácil para utilizar e manusear. Ele pode ser calibrado por meio do software gratuito o Tobii Pro Eyetracker Manager. Por ser um modelo mais simples, também é amplamente utilizado pelos streamers na transmissão de gameplays.
 - Segundo o site da Tobii, a tolerância de movimento da cabeça é boa, mantendo um bom nível de precisão, porém é mais sensível do que sistemas que utilizam o sistema de câmera dupla (este utiliza o sistema de câmera única). A distância de operação do eyetracker é de 45cm a 95cm.
- Psychopy: É um pacote open source construído por pesquisadores para pesquisadores, com o objetivo de facilitar a produção de rotinas de experimentos rotineiramente utilizadas em projetos nas áreas de psicolinguística, psicologia, neurociência, comportamento, entre outros. A plataforma é construída na linguagem Python, e permite que a adição de scripts na linguagem por parte do usuário, isso trás flexibilidade na plataforma e liberdade para o pesquisador produzir o seu próprio experimento. É uma alternativa gratuita para o MatLab. (Peirce et al., 2019)

Pela sua natureza volátil, o Psychopy permite a adição de scripts python ou Javascript em sua rotina. O software foi desenvolvido de tal forma que quando um novo experimento é criado, paralelamente um script em python é construído e é disponibilizado para o pesquisador alterar caso queira. Isso ajuda àqueles que já possuem domínio com programação e desejam melhorar o seu experimento. Desse modo, o software disponibiliza um componente chamado “Code Component”, que possui várias abas, essas que adicionam o código digitado pelo desenvolvedor em determinadas partes do algoritmo produzido automaticamente pelo experimento. Essas abas são: BEGIN EXPERIMENT, BEGIN ROUTINE, EVERY FRAME, END ROUTINE e END EXPERIMENT. Neste projeto utilizamos esse componente para adicionar novas funcionalidades no projeto. Além disso, mais uma vez pela forma como foi construído, o Psychopy permite a adição de variáveis em sua estruturação. Logo, é permitido que o pesquisador crie uma variável e a utilize nos componentes disponibilizados pelo próprio programa. No Psychopy, essas variáveis são representadas pelo uso do \$.

- PIE Context Data: Um banco de dados produzido pelo pesquisador original da pesquisa, contendo diversas MWEs (Multi-Word Expressions, Expressões de múltiplas palavras, em português), retiradas de diferentes fontes, como o MAGPIE e o SemEval22. Cada expressão idiomática contém duas frases exemplo, uma frase que reproduz o seu sentido idiomático e outra que reproduz o seu sentido literal. Além disso, como já foi mencionado anteriormente, o experimento possui expressões a

serem utilizadas na fase de teste, essas estão referenciadas por uma coluna (\$block) no banco de dados. Ainda mais, o banco de dados também contém as posições iniciais e finais (indexes) das expressões idiomáticas na frase.

- Comprehension Questions: As perguntas de compreensão estão em uma lista, sendo classificadas entre Numbers e Scalars.

Resultados e Discussões

O avanço deste projeto atingiu um marco fundamental com a aprovação do protocolo de pesquisa pelo Comitê de Ética em Pesquisa (CEP). A obtenção deste parecer favorável não apenas legitima os procedimentos metodológicos propostos, mas também assegura que a condução do estudo seguirá os mais rigorosos padrões éticos, garantindo a proteção e o bem-estar dos participantes. Esta etapa, agora superada, nos permitiu avançar para a fase de execução, que se encontra em pleno andamento com o recrutamento dos participantes e o início do processo de coleta de dados. Esta transição da fase de planejamento para a prática é crucial e representa o momento em que as hipóteses do estudo começarão a ser efetivamente investigadas.

Um desenvolvimento de particular relevância para o fortalecimento e a internacionalização da pesquisa foi a viagem técnica realizada para a Inglaterra, uma iniciativa que contou com o prestigioso apoio financeiro da Royal Society do Reino Unido. Este fomento não só viabilizou a imersão em um ambiente acadêmico de excelência, mas também chancelou a relevância e o potencial da nossa investigação perante uma das mais antigas e respeitadas sociedades científicas do mundo.

O ponto alto desta viagem foi a visita ao Alan Turing Institute, em Londres, o instituto nacional do Reino Unido para ciência de dados e inteligência artificial. O intercâmbio de conhecimentos e a observação das pesquisas de ponta desenvolvidas no instituto nos proporcionaram valiosos insights, especialmente no que tange às técnicas avançadas de análise de dados que planejamos empregar. O diálogo com pesquisadores e a exposição a novas abordagens metodológicas serviram para refinar nossa estratégia analítica e abriram perspectivas para futuras colaborações. Esta experiência, portanto, não foi apenas complementar, mas sim um elemento enriquecedor que agrega uma camada de sofisticação e robustez ao projeto.

Conclusão

O projeto encontra-se em uma fase promissora, alicerçado por três pilares essenciais: a validação ética, o início da coleta de dados e o fortalecimento de laços internacionais. A aprovação pelo Comitê de Ética confere a solidez necessária para uma pesquisa responsável e de credibilidade. O andamento da fase de recrutamento e coleta de dados marca o progresso tangível em direção aos nossos objetivos, nos aproximando das respostas que buscamos.

Adicionalmente, o apoio da Royal Society e a imersão no ecossistema de inovação do Alan Turing Institute foram cruciais para elevar o patamar do projeto, agregando conhecimento de vanguarda e reconhecimento internacional. Com uma base metodológica eticamente aprovada, um plano de coleta de dados em execução e uma perspectiva analítica enriquecida

pela experiência internacional, as expectativas para a conclusão bem-sucedida deste estudo são extremamente positivas. Os resultados a serem obtidos têm o potencial de gerar contribuições significativas para o campo, refletindo uma pesquisa conduzida com rigor, responsabilidade e colaboração global.

Referências

MI, M.; VILLAVICENCIO, A.; MOOSAVI, N. S. Rolling the DICE on Idiomaticity: How LLMs Fail to Grasp Context. Disponível em: <<https://arxiv.org/abs/2410.16069>>. Acesso em: 12 mar. 2025.

PHELPS, D. et al. Sign of the Times: Evaluating the use of Large Language Models for Idiomaticity Detection. ACL Anthology, p. 178–187, maio 2024.

HOLLENSTEIN, N.; BARRETT, M.; BJÖRNSDÓTTIR, M. The Copenhagen Corpus of Eye Tracking Recordings from Natural Reading of Danish Texts. ACL Anthology, p. 1712–1720, jun. 2022.

HAAGSMA, H.; BOS, J.; NISSIM, M. MAGPIE: A Large Corpus of Potentially Idiomatic Expressions. ACL Anthology, p. 279–287, maio 2020.

RAYNER, K. Eye movements and attention in reading, scene perception, and visual search. Quarterly Journal of Experimental Psychology, v. 62, n. 8, p. 1457–1506, ago. 2009..

ECE TAKMAZ et al. Generating Image Descriptions via Sequential Cross-Modal Alignment Guided by Human Gaze. Data Archiving and Networked Services (DANS), 1 jan. 2020.

DOS, C. combinação de palavras com sentido figurado. Disponível em: <https://pt.wikipedia.org/wiki/Express%C3%A3o_idiom%C3%A1tica>. Acesso em: 26 nov. 2025.

PEIRCE, J. et al. PsychoPy2: Experiments in behavior made easy. Behavior Research Methods, v. 51, n. 1, p. 195–203, fev. 2019.

RAYNER, K. et al. Eye movements during information processing tasks: Individual differences and cultural effects. Vision Research, v. 47, n. 21, p. 2714–2726, set. 2007.

CÓDIGO: ET1227

AUTOR: MATHEUS SENAS DE CRISTO

COAUTOR: ISAAC MARLON DA SILVA LOURENÇO

COAUTOR: NATALYA GABRIELE NUNES DE AZEVEDO

COAUTOR: JOSE EDUARDO MONTEIRO DOS SANTOS

ORIENTADOR: PATRICK CESAR ALVES TERREMATTE

TÍTULO: O desenvolvimento colaborativo de uma Plataforma Inteligente para Gestão Curricular e Apoio Acadêmico no Bacharelado em Tecnologia da Informação

Resumo

O projeto “Bora Pagar” propõe uma plataforma de planejamento curricular voltada aos alunos do Bacharelado em Tecnologia da Informação (BTI) do Instituto Metrópole Digital (IMD), visando solucionar dificuldades na escolha de disciplinas em um curso de matriz curricular flexível. Devido à falta de orientação eficaz e à escassez de professores, muitos alunos enfrentam problemas no planejamento acadêmico, o que resulta em sobrecarga, indeferimento de matrículas e organização desordenada. A plataforma pretende oferecer funcionalidades como visualização de interesse por disciplinas, simulação de matrículas e recomendações personalizadas, além de ser uma fonte de dados relevante para o planejamento acadêmico da secretaria do BTI, promovendo uma gestão mais estratégica e alinhada com as necessidades estudantis.

Palavras-chave: planejamento curricular, gestão acadêmica.

TITLE: The collaborative development of an Intelligent Platform for Curricular Management and Academic Support in the Bachelor's Degree in Information Technology

Abstract

The "Bora Pagar" project proposes a curriculum planning platform for students of the Bachelor's in Information Technology (BTI) program at the Metrópole Digital Institute (IMD), aiming to solve difficulties in choosing courses within a flexible curriculum structure. Due to a lack of effective guidance and a shortage of professors, many students face problems with their academic planning, which results in being overloaded, having their enrollments denied, and disorganized scheduling. The platform intends to offer features such as visualizing student interest in specific courses, enrollment simulations, and personalized recommendations. Furthermore, it will serve as a relevant data source for the academic planning of the BTI's administrative office, promoting more strategic management that is aligned with student needs.

Keywords: Curriculum planning, Academic management

Introdução

O Bacharelado em Tecnologia da Informação (BTI), oferecido pelo Instituto Metrópole Digital (IMD), caracteriza-se por sua estrutura curricular flexível, permitindo que os discentes escolham, além dos componentes curriculares obrigatórios, uma ampla variedade de disciplinas optativas. Embora essa liberdade proporcione uma formação diversificada e adaptável, também gera dificuldades para os alunos, que frequentemente encontram obstáculos ao tentar selecionar as disciplinas mais adequadas aos seus interesses e necessidades acadêmicas. Esse cenário evidencia um problema de orientação acadêmica, pois a quantidade de opções acaba por confundir muitos alunos, dificultando o planejamento de sua trajetória educacional. O grêmio estudantil, em conjunto com a secretaria do BTI, adota uma estratégia paliativa ao disponibilizar formulários para que os estudantes registrem seu interesse nas disciplinas de semestres futuros. No entanto, essa metodologia é limitada e imprecisa. Por não proporcionar um planejamento eficaz, os alunos tendem a demonstrar interesse em um número excessivo de disciplinas, na tentativa de aumentar suas chances de matrícula. Isso gera um planejamento desordenado e insuficiente, uma vez que o formulário, pela sua estrutura genérica, não permite uma análise detalhada das reais necessidades e intenções dos discentes. Como consequência, a secretaria enfrenta dificuldades para alinhar a oferta de disciplinas com a demanda efetiva dos alunos, o que prejudica o planejamento acadêmico e compromete a organização curricular dos próximos semestres. O planejamento da grade curricular é um fator essencial para o sucesso acadêmico e profissional dos estudantes. Um planejamento estruturado e balanceado permite que o aluno distribua suas disciplinas de maneira equilibrada ao longo do curso, otimizando o tempo e o esforço dedicados ao aprendizado. Além disso, uma organização curricular eficiente permite ao estudante identificar pré-requisitos, priorizar disciplinas essenciais e evitar sobrecargas em determinados períodos, o que aumenta as chances de sucesso e conclusão do curso. Em vista da dificuldade de obter esse equilíbrio, muitos alunos recorrem a métodos manuais de organização, como esboços em cadernos, quadros em plataformas digitais e planilhas eletrônicas. Entretanto, esses métodos são individuais e não permitem uma visão agregada das demandas dos alunos, dificultando a atuação da secretaria em resposta às necessidades específicas de cada discente. Neste contexto, foi desenvolvido o projeto “Bora Pagar”, uma plataforma integrada de gestão e planejamento curricular para os alunos do BTI, acessível em <https://matriculaai.imd.ufrn.br/>. A proposta envolve não apenas uma ferramenta de organização de matérias, mas também a implementação de funcionalidades que agregam valor ao planejamento acadêmico. A plataforma permite que os alunos visualizem a quantidade de interessados em cada disciplina, simulem pré-resultados de matrícula com base nos índices acadêmicos e obtenham sugestões de disciplinas de acordo com a demanda e o histórico individual. Além disso, a ferramenta serve como um recurso valioso para a secretaria do BTI, possibilitando um planejamento mais estratégico e alinhado com a demanda efetiva dos estudantes. Esse projeto é relevante para o campo da Tecnologia da Informação, pois trata de um problema significativo e recorrente em cursos de caráter flexível, contribuindo para o avanço do conhecimento sobre métodos e tecnologias que podem facilitar o gerenciamento acadêmico. Ao fomentar uma organização mais eficiente da vida acadêmica dos alunos, o “Bora Pagar” almeja impactar positivamente o desempenho dos discentes e a gestão de recursos educacionais, beneficiando diretamente o desenvolvimento acadêmico e a formação de profissionais mais preparados para o mercado.

Metodologia

O projeto será conduzido seguindo a metodologia ágil Scrum, garantindo uma abordagem interativa e incremental que facilita a adaptação às mudanças e melhorias contínuas. Para a

implementação técnica do projeto, foi selecionado um conjunto de tecnologias modernas e consolidadas no mercado, visando garantir a robustez, escalabilidade e manutenibilidade da solução. O backend será desenvolvido em Java 21, utilizando o framework Spring Boot para a construção de APIs RESTful. A segurança e autenticação dos usuários serão implementadas via OAuth 2.0, integrando a plataforma aos sistemas SIGs da instituição para permitir o Single Sign-On (SSO). O armazenamento e a persistência de dados ficarão a cargo do banco de dados PostgreSQL, com o controle de versionamento de esquema gerenciado pela ferramenta Flyway. No frontend, a interface do usuário será uma Single-Page Application (SPA) reativa, construída com Vue.js 3 e gerenciada pela ferramenta de build Vite. A estilização será feita com o framework Tailwind CSS. Para a gestão do projeto, o controle de versão e a prototipagem de interfaces, a equipe utilizará, respectivamente, as ferramentas GitHub Projects, Git e Figma. Abaixo está o fluxo adaptado às necessidades do projeto: 1. Planejamento de Sprint Objetivo: Definir funcionalidades chave para o planejamento curricular, simulação de matrículas, e integração com o sistema acadêmico, dividindo-as em tarefas específicas. Ação: Estabelecer prioridades de desenvolvimento, alinhando as atividades do backend, frontend e BI para entregar incrementos funcionais a cada sprint. Desenvolvimento Objetivo: Desenvolver as funcionalidades previstas de forma contínua, como a autenticação de alunos, APIs para planejamento curricular, visualizações de BI e interface de usuário. 2. Ação: Backend: Construir APIs para suportar a gestão curricular e integração com os sistemas acadêmicos. Frontend: Implementar interfaces intuitivas para o gerenciamento de planejamento e simulação de matrícula, com interação fluida entre backend e frontend. BI: Desenvolver funcionalidades de coleta e visualização de dados acadêmicos para análises estratégicas, ajudando na tomada de decisão sobre a oferta de disciplinas. Revisão Objetivo: Avaliar as entregas realizadas no sprint, garantindo que as funcionalidades atendem aos requisitos e objetivos definidos. Ação: Testar a integração entre backend e frontend, revisar a usabilidade das interfaces e verificar a precisão dos dados apresentados nas análises de BI.

Resultados e Discussões

O percurso acadêmico no Bacharelado em Tecnologia da Informação (BTI) é, por definição, uma jornada de autogestão. A sua matriz curricular, aberta e flexível, oferece aos alunos uma liberdade ímpar para moldar a sua formação, mas, paradoxalmente, essa mesma liberdade pode gerar um ambiente de incerteza e complexidade. A vivência dos alunos é marcada pela dificuldade em montar uma grade de horários compatíveis e pela incerteza sobre o alinhamento das disciplinas com a carreira, além do risco constante de descobrir pendências de pré-requisitos em cima da hora da matrícula. 1. A Origem Colaborativa: Uma Solução Nascida da Necessidade O projeto “Bora Pagar” não nasceu de uma demanda institucional, mas de algo muito mais genuíno: uma dor compartilhada entre colegas. A iniciativa partiu de um esforço colaborativo entre nós, discentes, que, ao vivenciarmos os mesmos problemas, muitos deles intensificados pela crescente complexidade dos currículos interdisciplinares na educação superior (RAMOS, 2016), nos unimos para criar uma solução que pudesse beneficiar toda a nossa comunidade acadêmica. Atuando no desenvolvimento full stack (backend e frontend) deste projeto, minhas principais contribuições focaram-se na implementação de funcionalidades que buscassem, acima de tudo, trazer tranquilidade e clareza para os estudantes 2. A Arquitetura da Confiança: Modelagem para Integração com o SIGAA O pilar técnico mais transformador do projeto “Bora Pagar” foi a capacidade de integrar os dados fornecidos pelo serviço do SIGAA. Meu trabalho principal consistiu em modelar e construir a arquitetura do sistema para que ele pudesse receber, interpretar e

organizar essa informação de forma inteligente, transformando-a em um perfil acadêmico coeso com o qual o aluno pudesse interagir. O desafio central foi projetar as entidades do nosso sistema, criando um espelho estruturado e compreensível dos dados do SIGAA. Essa tarefa foi agravada por uma característica desafiadora da API: a dificuldade em obter dados completos com uma única requisição. Muitas vezes, era necessário realizar múltiplas chamadas a diferentes endpoints e combinar suas respostas para formar uma única informação coesa, o que introduz complexidade e potenciais gargalos de performance. Diante desse cenário, uma decisão arquitetônica crucial foi adotar uma abordagem híbrida. Após um estudo detalhado que analisou a frequência de requisições, a volatilidade dos dados e o tempo de resposta da API, optou-se pela persistência de dados estratégicos. Essa solução de modelagem e persistência foi a fundação que permitiu mitigar a latência da API e transformar uma massa de informações complexas em um perfil de usuário funcional e confiável. A arquitetura deste fluxo pode ser visualizada na Figura 1, que destaca a contribuição principal na modelagem do sistema.

3. O Sistema de Interesses: Transformando Planejamento em Dados A funcionalidade mais estratégica, e uma das minhas principais contribuições, foi o desenvolvimento do sistema de manifestação de interesses. No frontend, desenvolvi a interface que permite ao aluno, de forma simples e intuitiva, selecionar as disciplinas que deseja cursar. No backend, fui responsável por modelar e construir toda a lógica para registrar e gerenciar essas intenções. Este sistema é a ponte direta entre o planejamento individual do aluno e a inteligência de dados que será disponibilizada para a secretaria.

4. A Rede Social: Planejando Juntos Partindo do desafio de que o planejamento da trajetória acadêmica é muitas vezes uma jornada solitária, concebi e desenvolvi o "Bora pagar", um ecossistema social completo. Para materializar essa visão, dediquei um esforço considerável no desenvolvimento de toda a arquitetura do sistema. Colaborei com a construção da solução completa, desde a modelagem do banco de dados e a lógica de backend, até a construção das telas de interação no frontend. A funcionalidade principal criada, a qual permite aos alunos se conectarem com base em interesses curriculares comuns, é a aplicação prática dos princípios de integração estudantil. Segundo Tinto (1993), a integração social não é um mero complemento, mas um pilar para o sucesso e a permanência estudantil. Dessa forma, o "Bora pagar" transcende sua função como ferramenta de organização; ele atua como um catalisador para a formação do capital social, fortalecendo a rede de apoio que mitiga o isolamento e enriquece a experiência universitária em sua totalidade.

Discussão: Unindo Pessoas, Processos e Tecnologia O "Bora Pagar" transcende a definição de uma simples ferramenta. Ele representa, em sua essência, um caso de sucesso de desenvolvimento colaborativo, onde os próprios usuários do sistema são seus criadores. Para garantir a qualidade em um ambiente com múltiplos desenvolvedores, estabelecemos um processo de trabalho organizado, com revisão de código em pares, verificações automáticas e reuniões periódicas, que eram o verdadeiro motor do nosso progresso. Essa disciplina técnica foi guiada por referências consolidadas da Engenharia de Software. No desenvolvimento do backend, por exemplo, buscamos aplicar os princípios de simplicidade e manutenibilidade de acordo com Martin (2020). De forma análoga, para o frontend, a construção de interfaces reativas foi inspirada nas boas práticas do ecossistema Vue.js (FREITAS, 2020). O projeto é, portanto, um caso de estudo sobre como a colaboração estudantil, quando organizada com processos de engenharia e embasada na literatura da área, pode produzir soluções de alto impacto, focadas em resolver problemas reais de pessoas reais.

Conclusão

Uma pesquisa de mercado aprofundada revelou um desafio central na gestão acadêmica das universidades públicas brasileiras: a ausência de um ecossistema que integre as necessidades de planejamento dos estudantes com a otimização da oferta de disciplinas pela instituição. O que existe hoje é uma lacuna entre a ansiedade discente, gerada pela incerteza da matrícula, e a capacidade administrativa de responder a essa demanda de forma estratégica. O projeto "Bora Pagar" foi concebido como uma resposta direta a essa necessidade, posicionando-se como uma solução inovadora para o cenário nacional. Neste trabalho, minha responsabilidade foi a implementação full stack das funcionalidades centrais, traduzindo as dores da comunidade estudantil em uma plataforma robusta e intuitiva. A conclusão fundamental deste trabalho é que a automação e a modelagem inteligente de dados são os pilares para um planejamento acadêmico mais estratégico e, acima de tudo, mais humano. O marco técnico mais significativo da minha contribuição foi o desenvolvimento de uma arquitetura de sistema capaz de interpretar e dar sentido à complexa estrutura de informações do SIGAA. As funcionalidades já entregues, como o sistema de manifestação de interesses e a rede social acadêmica, representam um avanço concreto, oferecendo ao estudante um nível inédito de segurança e controle sobre sua trajetória, o que contribui diretamente para a redução da ansiedade associada a este processo. Com base nos resultados sólidos alcançados, a evolução do projeto está planejada em três frentes estratégicas: ● Desenvolvimento Preditivo para o Aluno: Implementar um algoritmo de simulação de matrícula, uma ferramenta desenhada para oferecer aos estudantes a confiança necessária para tomarem as melhores decisões para seu futuro acadêmico. ● Inteligência de Dados para a Gestão: Desenvolver painéis de Business Intelligence para a coordenação de curso, completando a visão do projeto de ser uma ponte de colaboração que beneficia toda a comunidade acadêmica. ● Aprimoramento da Experiência Humana: Conduzir estudos de usabilidade aprofundados, pois compreendemos que o feedback de quem utiliza a ferramenta é o nosso mais valioso guia para a evolução contínua. Em síntese, minhas contribuições para o projeto "Bora Pagar" resultaram em uma base tecnológica sólida e de alto impacto. O trabalho realizado não apenas entregou funcionalidades essenciais, mas também estabeleceu um roteiro claro para o futuro, com o potencial de transformar positivamente a jornada dos estudantes e otimizar a gestão de recursos na instituição.

Referências

- RAMOS, Luiza Olivia Lacerda. O lugar da interdisciplinaridade na educação superior: uma análise dos projetos pedagógicos dos Cursos de Bacharelado Interdisciplinar da UFBA. 2016. Tese (Doutorado em Educação) - Universidade Federal da Bahia, Salvador, 2016. Disponível em: <https://repositorio.ufba.br/handle/ri/19621>. Acesso em: 6 nov. 2024.
- MARTIN, Robert Cecil. Código Limpo: Habilidades Práticas do Agile Software. São Paulo: Alta Books, 2019.
- MARTIN, Robert Cecil. Arquitetura Limpa: O Guia do Artesão para Estrutura e Design de Software. São Paulo: Alta Books, 2020.
- ALVES, Leandro Daniel. Padrões de Projeto: Soluções Reutilizáveis de Software Orientado a Objetos. São Paulo: Novatec Editora, 2016.
- FREITAS, Wilson; NETO, Marcelo Andrade. Aplicações Modernas com Vue.js. São Paulo: Casa do Código, 2020.
- TINTO, Vincent. Leaving college: rethinking the causes and cures of student attrition. 2. ed. Chicago: The University of Chicago Press, 1993.

Anexos

Figura 1: Fluxo de dados do SIGAA para o modelo de entidades do Bora Pagar

Figura 2: Tela principal do Sistema

CÓDIGO: ET1350

AUTOR: ANA CATARINA GONÇALVES FUENTES

COAUTOR: LORENZO SANTOS FURTADO

ORIENTADOR: RENAN CIPRIANO MOIOLI

TÍTULO: Desenvolvimento de Arquiteturas de Redes Neurais Pulsadas para Processamento de Sinais de Eletrocardiograma (ECG)

Resumo

Este trabalho investigou a viabilidade de execução de Redes Neurais Pulsantes (SNNs) em ambientes de supercomputação, tendo como objetivo final sua aplicação no processamento de sinais eletrocardiográficos (ECG). Contudo, devido às limitações tecnológicas encontradas — como alta demanda computacional, dependência de GPUs e estágio de desenvolvimento das bibliotecas disponíveis — tornou-se necessário iniciar com a classificação de imagens médicas, por serem mais acessíveis e de manipulação menos complexa que séries temporais. Embora o Google Colab tenha se mostrado adequado para a prototipagem inicial, sua restrição de memória e processamento inviabilizou o treinamento em larga escala, especialmente com o dataset ISIC de câncer de pele. Assim, foram conduzidos esforços de adaptação de bibliotecas como PyGeNN e mlGeNN ao Núcleo de Processamento de Alto Desempenho (NPAD/UFRN), incluindo análise de compatibilidade, configuração de dependências e adequação ao uso de GPUs. Os resultados preliminares indicaram que, em ambientes limitados, o treinamento de SNNs aplicado a dados reais resulta em baixa acurácia e elevado custo computacional. Entretanto, a infraestrutura de alto desempenho representa um caminho promissor para viabilizar aplicações práticas das SNNs em cenários biomédicos, especialmente no processamento de ECG.

Palavras-chave: Redes Neurais Pulsadas; Supercomputação; Processamento de Imagens Médi

TITLE: Development of Spiking Neural Network Architectures for Electrocardiogram Signal Processing

Abstract

This study investigated the feasibility of running Spiking Neural Networks (SNNs) in high-performance computing environments, with the ultimate goal of applying them to electrocardiogram (ECG) signal processing. However, due to technological limitations — such as high computational demand, dependence on GPUs, and the early development stage of the available libraries — it was necessary to begin with the classification of medical images, which are more accessible and less complex to handle than time series. Although Google Colab proved adequate for initial prototyping, its memory and processing restrictions made large-scale training unfeasible, especially with the ISIC skin cancer dataset. Consequently, efforts were directed toward adapting libraries such as PyGeNN and mlGeNN to the High-Performance Computing Center (NPAD/UFRN), including compatibility

analysis, dependency configuration, and adjustment for GPU usage. Preliminary results indicated that, in limited environments, training SNNs on real data leads to low accuracy and high computational cost. Nonetheless, high-performance infrastructure represents a promising path to enable practical applications of SNNs in biomedical scenarios, particularly for ECG signal processing.

Keywords: Spiking Neural Networks; High-Performance Computing; Medical Image Pro

Introdução

O avanço das redes neurais artificiais têm impulsionado o desenvolvimento de diversos métodos de aprendizado profundo, com aplicações que abrangem desde a visão computacional até a biomedicina. Entre essas abordagens, uma corrente neuromórfica surge: as redes neurais pulsantes (Spiking Neural Networks – SNNs). Essa abordagem bioinspirada vem ganhando destaque por se aproximar do funcionamento biológico dos neurônios, apresentando potencial para maior eficiência energética e novos paradigmas de processamento de informação.

Apesar desse potencial, o treinamento de SNNs ainda enfrenta limitações significativas. Um dos principais desafios é a elevada demanda por recursos de hardware, especialmente unidades de processamento gráfico (GPU), que restringe a realização de experimentos em ambientes acadêmicos com infraestrutura limitada. Além disso, as bibliotecas que implementam esse tipo de rede encontram-se em estágios iniciais de desenvolvimento, carecendo de documentação abrangente e apresentando comportamento pouco transparente, o que dificulta a adaptação a diferentes plataformas de computação.

Nesse contexto, o objetivo central deste trabalho é o processamento de sinais eletrocardiográficos (ECG) utilizando redes neurais pulsantes. Entretanto, as limitações tecnológicas encontradas exigiram a realização de uma etapa prévia de adaptação e preparação, tanto para viabilizar a execução dos experimentos em ambientes de supercomputação quanto para explorar inicialmente o uso de imagens como dado de entrada. Essa escolha se justifica pelo fato de que imagens, como no caso de bases médicas já consolidadas, oferecem maior acessibilidade experimental e simplicidade de manipulação em comparação às séries temporais do ECG. Assim, o estudo busca não apenas avaliar os obstáculos técnicos e computacionais para a aplicação das SNNs em contextos biomédicos, mas também estabelecer caminhos para sua futura aplicação em sinais fisiológicos mais complexos.

Metodologia

A metodologia adotada neste trabalho foi organizada em quatro etapas principais, com o objetivo de avaliar o desempenho das redes neurais pulsadas (SNNs) e investigar os desafios técnicos de sua aplicação em diferentes contextos. Inicialmente, foram conduzidos experimentos em conjuntos de dados de referência amplamente utilizados na literatura, como o MNIST e o DermaMNIST, de modo a verificar a correta implementação das bibliotecas PyGeNN e mlGeNN e estabelecer uma linha de base de desempenho em cenários já consolidados. Em seguida, buscou-se aproximar o estudo de situações mais próximas da prática médica, aplicando os modelos ao conjunto de dados ISIC, que reúne imagens reais de

câncer de pele. Essa etapa permitiu identificar dificuldades adicionais relacionadas à maior variabilidade e complexidade das imagens, evidenciando as limitações do ambiente de execução em nuvem.

A partir dos resultados obtidos nessa fase, foi realizada uma análise da demanda computacional necessária para o treinamento das SNNs em problemas de maior escala. Observou-se que a execução em GPU é indispensável e que os recursos oferecidos pelo Google Colab não eram suficientes para suportar o processamento de imagens em alta resolução e por longos períodos de treinamento. Essa constatação reforçou a necessidade de explorar soluções de maior capacidade, especialmente em ambientes de supercomputação.

Dessa forma, a última etapa consistiu na adaptação das bibliotecas e dos pacotes utilizados, de modo a viabilizar sua execução no Núcleo de Processamento de Alto Desempenho (NPAD/UFRN). Esse processo envolveu a análise de dependências, a verificação da compatibilidade das versões do CUDA e a reestruturação do código para submissão em jobs via SLURM. Embora a execução definitiva ainda não tenha sido concluída devido a dificuldades de reconhecimento da GPU pelo framework, essa preparação foi essencial para diagnosticar os principais obstáculos técnicos e estabelecer os requisitos necessários para a continuidade da pesquisa em ambiente de alto desempenho.

Resultados e Discussões

Durante a execução em ambiente de nuvem, utilizando o Google Colab, observou-se que as limitações de memória e processamento impactaram significativamente o desempenho das redes neurais pulsadas. Para o conjunto de dados MNIST, a rede foi treinada por 10 épocas, utilizando 256 neurônios de entrada (um para cada pixel da imagem 28×28), uma camada escondida com 128 neurônios Leaky Integrate Fire (LIF) e 10 neurônios de saída Leaky Integrate (LI) com conectividade densa. A performance, avaliada com SparseCategoricalAccuracy, apresentou acurácia de 0,9317 no conjunto de treino e 0,9372 no conjunto de teste, com tempo de treinamento em torno de 25 minutos.

Na sequência, os mesmos hiperparâmetros foram aplicados ao DermaMNIST, resultando em acurácia de 0,59 após 10 épocas de treinamento e 0,63 quando treinada por 40 épocas, com tempo de execução aproximado de 1 hora e 30 minutos. Esses resultados demonstram que, embora a arquitetura funcione em conjuntos de referência, a complexidade e a variabilidade dos dados médicos exigem maior quantidade de iterações e ajustes de hiperparâmetros para atingir desempenho satisfatório.

No caso do conjunto ISIC, devido à alta resolução das imagens (600×450 pixels), foi necessário reduzir o tamanho para 120×90 pixels para contornar falhas recorrentes por esgotamento de memória. A melhor configuração exploratória foi obtida com 20 épocas, $\Theta = 0,81$ e 758 neurônios na camada escondida, alcançando 65% de acurácia no décimo experimento de treino, com tempo de treinamento em torno de 45 minutos. Esses resultados evidenciam que a execução em Colab ainda não é suficiente para problemas biomédicos de grande escala e que ajustes finos de hiperparâmetros, juntamente com aumento de recursos computacionais, são fundamentais.

Em função dessas limitações, iniciou-se a adaptação da implementação para o ambiente de supercomputação do NPAD, incluindo a instalação das bibliotecas PyGeNN e mlGeNN,

configuração de dependências e preparação do código para submissão em jobs via SLURM. No entanto, surgiram obstáculos relacionados à compatibilidade com o CUDA disponível no cluster, de modo que a biblioteca ainda não consegue reconhecer corretamente a GPU, impossibilitando a execução do treinamento. Apesar disso, o trabalho avançou na preparação do ambiente e no diagnóstico dos problemas técnicos, permitindo compreender com maior clareza as barreiras que precisam ser superadas para viabilizar experimentos robustos em HPC.

Esses resultados parciais destacam duas constatações principais. Primeiro, que a aplicação de SNNs em bases biomédicas complexas é inviável em plataformas de recursos limitados, dada a elevada demanda computacional. Segundo, que a migração para ambientes de supercomputação é necessária, mas ainda depende da maturidade das bibliotecas e de sua compatibilidade com drivers e GPUs modernas. Assim, o trabalho evidencia tanto o potencial das SNNs quanto os desafios práticos de sua implementação em larga escala, indicando a necessidade de avanços técnicos para consolidar essa linha de pesquisa em cenários biomédicos reais.

Conclusão

Este estudo demonstrou que a aplicação prática de Redes Neurais Pulsadas em bases médicas reais ainda enfrenta obstáculos significativos quando realizada em ambientes de recursos limitados. O Google Colab mostrou-se adequado para testes iniciais, mas insuficiente para treinar redes em larga escala.

A adaptação das bibliotecas PyGeNN e mlGeNN ao ambiente do NPAD/UFRN representa um passo essencial para viabilizar a pesquisa em supercomputação, permitindo explorar resoluções mais altas, tempos de treinamento reduzidos e maior estabilidade.

Como trabalhos futuros, pretende-se concluir a execução do dataset ISIC em HPC, investigar o uso de múltiplas GPUs em paralelo e comparar a eficiência energética das execuções em diferentes ambientes. Tais avanços são fundamentais para validar a viabilidade das SNNs em aplicações biomédicas e aproximar sua implementação do hardware neuromórfico.

Referências

YAVUZ, Esin; TURNER, James; NOWOTNY, Thomas. GeNN: a code generation framework for accelerated brain simulations. *Scientific Reports*, [S.L.], v. 6, n. 1, 7 Jan. 2016. Springer Science and Business Media LLC. <http://dx.doi.org/10.1038/srep18854>.

NVIDIA, Vingelmann, P., and Fitzek, F. H. (2020). CUDA, Developer. Available online at: <https://nvidia.com/cuda-toolkit>.

KNIGHT, James C.; KOMISSAROV, Anton; NOWOTNY, Thomas. PyGeNN: a python library for gpu-enhanced neural networks. *Frontiers In Neuroinformatics*, [S.L.], v. 15, 22 abr. 2021. Frontiers Media SA. <http://dx.doi.org/10.3389/fninf.2021.659005>.

KNIGHT, James C; NOWOTNY, Thomas. Easy and efficient spike-based Machine Learning with mlGeNN. Neuro-Inspired Computational Elements Conference, [S.L.], p. 115-120, 11 abr. 2023. ACM. <http://dx.doi.org/10.1145/3584954.3585001>.

NOWOTNY, Thomas; TURNER, James P; KNIGHT, James C. Loss shaping enhances exact gradient learning with Eventprop in spiking neural networks. Neuromorphic Computing And Engineering, [S.L.], v. 5, n. 1, p. 014001, 21 Jan. 2025. IOP Publishing. <http://dx.doi.org/10.1088/2634-4386/ada852>.

HABAEBI, Mohamed Hadi; ZULKARNAIN, Muhammad Amin; GUNAWAN, Teddy Surya; FITRIAWAN, Helmy; KARTIWI, Mira; ABD.RAHMAN, Faridah. Development of Skin Cancer Detection Using Deep Spiking Neural Networks. 2024 Ieee 10Th International Conference On Smart Instrumentation, Measurement And Applications (Icsima), [S.L.], p. 125-129, 30 jul. 2024. IEEE. <http://dx.doi.org/10.1109/icsima62563.2024.10675560>.

CÓDIGO: ET1462

AUTOR: HENRIQUE LOPES FOUQUET

ORIENTADOR: CHARLES ANDRYE GALVAO MADEIRA

TÍTULO: Desenvolvimento de módulos da plataforma computacional 2.0 para a experimentação de modelos de aprendizado profundo por reforço em supercomputação com GPUs

Resumo

Este trabalho apresenta a continuidade do desenvolvimento da versão 2.0 da plataforma URNAI, destinada à experimentação de modelos de aprendizado por reforço profundo em supercomputação com GPUs. Durante este período, o foco foi a integração da plataforma com a biblioteca *Stable Baselines3*, permitindo a utilização de algoritmos amplamente empregados pela comunidade científica, como *PPO* e *A2C*. Para viabilizar essa integração, foi criada a classe *CustomEnv*, responsável por conectar os módulos de Ambiente, Estado, Espaço de Ações e Recompensa da URNAI com os modelos do SB3. Outra contribuição relevante foi a integração com o *Weights & Biases (wandb)*, possibilitando o registro e a visualização em tempo real de métricas experimentais, além da comparação sistemática entre execuções. No momento, a equipe está conduzindo a reprodução dos experimentos da DeepMind em minigames de *StarCraft II*, como *CollectMineralShards*, *DefeatRoaches* e *BuildMarines*. Esses avanços reforçam o caráter modular e científico da URNAI 2.0, preparando a plataforma para experimentos mais robustos e comparáveis no futuro.

Palavras-chave: Aprendizado por
Reforço.Supercomputação.StableBaselines3.StarCraft II.

TITLE: Development of Modules for the Computational Platform 2.0 for the Experimentation of Deep Reinforcement Learning Models in Supercomputing with GPUs

Abstract

This work presents the continuation of the development of URNAI version 2.0, a platform designed for experimenting with deep reinforcement learning models in GPU-based supercomputing. During this period, the focus was on integrating the platform with the *Stable Baselines3* library, enabling the use of widely adopted algorithms such as *PPO* and *A2C*. To enable this integration, a *CustomEnv* class was developed, responsible for connecting URNAI's Environment, State, Action Space, and Reward modules to SB3 models. Another important contribution was the integration with *Weights & Biases (wandb)*, which allows real-time logging and visualization of experimental metrics, as well as systematic comparison between runs. At the current stage, the team is working on reproducing DeepMind's experiments in *StarCraft II* minigames, such as *CollectMineralShards*, *DefeatRoaches*, and *BuildMarines*. These advances highlight the modular and scientific nature of URNAI 2.0, preparing the platform for more robust and comparable experiments in the future.

Keywords: Reinforcement Learning. Supercomputing. StableBaselines3.StarCraft II.

Introdução

O avanço da Inteligência Artificial (IA) tem sido marcado por resultados expressivos em domínios de alta complexidade, sobretudo após o sucesso da DeepMind com o AlphaGo, que demonstrou a capacidade do aprendizado por reforço profundo (*Deep Reinforcement Learning - DRL*) em vencer campeões humanos no jogo de tabuleiro *Go* (Silver et al., 2015; Gibney, 2016). A combinação entre redes neurais profundas e aprendizado por reforço (Sutton & Barto, 2018; Goodfellow, 2016) mostrou-se uma técnica poderosa para problemas de tomada de decisão sequenciais. Após esse marco, a comunidade científica voltou sua atenção para ambientes ainda mais desafiadores, como o jogo de estratégia em tempo real *StarCraft II*, que apresenta características de ambientes multiagente, parcialmente observáveis, estocásticos e de grande dimensionalidade (Vinyals et al., 2017).

Dentro desse contexto, foi desenvolvido o projeto URNAI (UFRN Artificial Intelligence), uma plataforma modular para experimentação de modelos de DRL em diferentes ambientes, como *PySC2* (para *StarCraft II*), *OpenAI Gym*, *ViZDoom* e *DeepRTS* (Madeira et al., 2020). A primeira versão da URNAI possibilitou avanços importantes em termos de flexibilidade e modularidade, mas com o crescimento da complexidade dos algoritmos e dos cenários utilizados, surgiram limitações relacionadas à escalabilidade e à manutenção do código. Isso motivou a criação da URNAI 2.0, uma versão mais otimizada, com ênfase em eficiência, qualidade e capacidade de integração com ambientes de supercomputação baseados em GPUs.

O trabalho desenvolvido no ano anterior concentrou-se principalmente na refatoração do código e na implementação de testes unitários, visando garantir maior confiabilidade e robustez ao software. Nesta nova etapa, o foco recaiu sobre a integração da plataforma com a biblioteca *Stable Baselines3* (SB3), uma das mais utilizadas pela comunidade de aprendizado por reforço, e com a ferramenta *Weights & Biases (wandb)*, amplamente adotada para acompanhamento e visualização de experimentos em aprendizado de máquina. Essas integrações permitem tanto o uso de algoritmos de referência quanto o monitoramento detalhado do desempenho dos agentes, facilitando a análise e a comparação de resultados.

Atualmente, o projeto encontra-se em fase de experimentação com minigames disponibilizados pela DeepMind, como *CollectMineralShards*, *DefeatRoaches* e *BuildMarines*. O objetivo é reproduzir resultados já publicados e validar a confiabilidade da plataforma, estabelecendo assim uma base sólida para futuras investigações. Essa estratégia de reprodutibilidade não apenas fortalece a credibilidade científica da ferramenta, mas também assegura que os módulos desenvolvidos possam ser utilizados em pesquisas futuras de maneira transparente e comparável.

Dessa forma, o desenvolvimento da URNAI 2.0 representa uma contribuição relevante para a área de aprendizado por reforço profundo, ao aliar qualidade de software, integração com ferramentas consolidadas e reprodutibilidade experimental. A expectativa é que, com esses avanços, a plataforma se torne uma referência prática para a condução de experimentos em ambientes complexos como o *StarCraft II*, além de abrir caminho para a exploração de novos desafios e aplicações.

Metodologia

Este trabalho foi desenvolvido com base na abordagem de pesquisa-ação (Thiollent, 2011), adequada por aliar a resolução de problemas práticos à produção de conhecimento científico. O processo seguiu ciclos iterativos de planejamento, desenvolvimento, teste e avaliação, permitindo a evolução incremental da plataforma URNAI 2.0.

O desenvolvimento da ferramenta foi guiado pelos seguintes princípios metodológicos:

- **Modularidade da URNAI:** cada parte essencial de um cenário de Aprendizado por Reforço (Ambiente, Estado, Espaço de Ações e Recompensa) foi estruturada como um módulo independente, permitindo substituição e reconfiguração conforme a necessidade. Essa característica foi fundamental para possibilitar a integração planejada com bibliotecas externas, como a *Stable Baselines3* (SB3);
- **Integração com bibliotecas externas:** para tornar a plataforma compatível com algoritmos amplamente utilizados, foi planejada a criação de uma classe *CustomEnv*, seguindo o padrão do *Gym*. Essa classe organiza as funcionalidades já implementadas (Ambiente, Estado, Ações e Recompensa) e viabiliza a comunicação direta com modelos do SB3, como *PPO*;
- **Práticas de engenharia de software:** o versionamento de código foi realizado no *GitHub*, utilizando branches e issues para organização das tarefas. O processo de integração foi controlado por ferramentas de *linter*, garantindo que *merges* só fossem aceitos quando o código estivesse livre de erros;
- **Garantia de qualidade:** testes unitários foram implementados com *pytest* e *unittest*, de forma a validar o correto funcionamento de cada módulo de maneira isolada;
- **Ambiente padronizado:** o uso do *Docker* permitiu encapsular todas as dependências da plataforma, possibilitando sua execução em diferentes máquinas sem problemas de configuração. Isso contribuiu para a reprodutibilidade dos experimentos;
- **Ciclos incrementais de desenvolvimento:** a cada entrega parcial do software, novos testes foram planejados em cenários progressivamente mais complexos, utilizando minigames de *StarCraft II* como base experimental.

Assim, a metodologia adotada combina boas práticas de engenharia de software com estratégias de modularidade e integração incremental, fornecendo as bases necessárias para que a URNAI 2.0 evolua como uma plataforma robusta e flexível para experimentação em Aprendizado por Reforço Profundo.

Resultados e Discussões

Ao longo do período de desenvolvimento, a URNAI 2.0 consolidou avanços significativos tanto no aspecto de qualidade de software quanto na integração com ferramentas de uso amplo na comunidade de DRL. Esses avanços tornam a plataforma mais robusta, reprodutível e adaptável a diferentes cenários experimentais, ampliando seu potencial de contribuição para pesquisas em ambientes de alta complexidade como o *StarCraft II*.

Integração com a *Stable Baselines3* (SB3)

Um dos principais resultados obtidos foi a integração prática com a biblioteca *Stable Baselines3* (SB3), reconhecida por reunir implementações estáveis e bem documentadas de algoritmos de DRL, como *PPO*, *A2C*, *DQN*, entre outros. Para viabilizar essa integração, foi criada a classe *CustomEnv*, que segue o padrão do *Gym* e atua como ponto de comunicação

entre os módulos da URNAI (Ambiente, Estado, Espaço de Ações e Recompensa) e os modelos do SB3.

Essa arquitetura modular permitiu que os algoritmos da SB3 fossem aplicados diretamente aos ambientes controlados pela URNAI, sem necessidade de alterações profundas na estrutura já existente. Como resultado, tornou-se possível treinar agentes em minigames de *StarCraft II* utilizando algoritmos altamente reconhecidos pela comunidade científica, ampliando consideravelmente o alcance e a aplicabilidade da plataforma.

Suporte ao *Action Masking*

Outro avanço relevante foi a implementação de um método no Espaço de Ações capaz de identificar, a cada *step*, quais ações estão disponíveis ou não para o agente. Esse recurso, embora não configure a implementação completa do *Action Masking*, fornece informações que podem ser utilizadas por modelos do SB3 que já oferecem suporte nativo a essa técnica.

O *Action Masking* é particularmente importante em ambientes como o *StarCraft II*, onde o agente frequentemente se depara com ações inválidas dependendo do contexto (por exemplo, tentar construir uma unidade sem recursos suficientes). Ao evitar que o agente selecione essas ações, o processo de aprendizado torna-se mais eficiente, uma vez que o esforço de exploração não é desperdiçado em comandos impossíveis. Assim, a URNAI passa a oferecer uma funcionalidade essencial para treinos mais rápidos e consistentes.

Registro e Visualização de Experimentos com *Weights & Biases (wandb)*

Outro resultado relevante foi a integração da plataforma com o *Weights & Biases (wandb)*, ferramenta amplamente utilizada na comunidade de aprendizado de máquina para registro e acompanhamento de experimentos. Essa integração possibilitou que métricas como recompensas médias, perdas e evolução do agente ao longo do tempo fossem registradas de forma automática e visualizadas em gráficos dinâmicos e interativos.

O uso do *wandb* trouxe benefícios em diferentes aspectos:

- **Monitoramento em tempo real:** tornou possível acompanhar a evolução do agente durante o treinamento sem a necessidade de gerar relatórios manuais;
- **Comparação entre experimentos:** os resultados de diferentes execuções podem ser comparados facilmente, auxiliando na escolha de hiperparâmetros e estratégias mais eficazes;
- **Reprodutibilidade e colaboração:** os experimentos ficam salvos em um repositório online, permitindo que a equipe compartilhe os resultados e reproduza as mesmas condições em execuções futuras.

Essa integração elevou a qualidade do processo experimental, pois agora não apenas o código está bem estruturado e testado, mas também os resultados científicos podem ser analisados de forma organizada e transparente. Isso reforça o papel da URNAI 2.0 como uma plataforma voltada não apenas ao desenvolvimento de agentes, mas também à condução de experimentos rigorosos e comparáveis no contexto de *DRL*.

Adoção do *Docker* e Reprodutibilidade

Um resultado igualmente importante foi a utilização do *Docker* para encapsular a execução da plataforma. Essa solução trouxe padronização ao ambiente de desenvolvimento e execução, permitindo que o software fosse executado em diferentes máquinas e servidores sem problemas de configuração. No contexto da pesquisa científica, isso significa maior reprodutibilidade dos experimentos, aspecto fundamental para validação e comparação de resultados.

Antes do uso do *Docker*, a instalação de dependências e configurações específicas poderia comprometer a continuidade do desenvolvimento. Agora, qualquer colaborador pode reproduzir exatamente o mesmo ambiente experimental, garantindo consistência ao longo do tempo e reduzindo falhas causadas por diferenças de infraestrutura.

Garantia de Qualidade do Código

A manutenção da qualidade do software continuou sendo uma prioridade no desenvolvimento da URNAI 2.0. O processo de versionamento de código no *GitHub* foi organizado com branches e issues, permitindo melhor controle das atividades. Além disso, o uso de *linters* garantiu que apenas código livre de erros fosse integrado, evitando que falhas se propagassem ao longo do desenvolvimento.

Os testes unitários, desenvolvidos com *pytest* e *unittest*, asseguraram que cada módulo funcionasse corretamente de forma isolada. A alta cobertura de testes trouxe confiabilidade ao processo, reduzindo significativamente o risco de erros durante a integração.

Experimentos com Minigames de *StarCraft II*

No campo experimental, os esforços foram direcionados para a reprodução dos resultados da DeepMind em minigames de *StarCraft II*, como *CollectMineralShards*, *DefeatRoaches* e *BuildMarines*. Cada um desses cenários apresenta desafios específicos:

- ***CollectMineralShards***: exige movimentação rápida e eficiente das unidades para coletar recursos dispersos, testando a capacidade do agente de planejar rotas;
- ***DefeatRoaches***: foca no combate direto, onde o agente precisa balancear ataque e sobrevivência;
- ***BuildMarines***: envolve a execução de uma sequência de ações de produção, exigindo coordenação e alocação de recursos.

Estes ambientes representam níveis progressivos de complexidade e servem como etapas de validação para a URNAI 2.0. Até o momento, a reprodução dos resultados da DeepMind encontra-se em andamento, mas já foi possível comprovar o funcionamento estável da integração com a SB3, o registro de métricas automáticas com o *wandb* e a adaptação da plataforma para uso com *Action Masking*.

Resultados Preliminares – *CollectMineralShards*

Nos experimentos realizados até o momento com o minigame *CollectMineralShards*, já é possível observar alguns avanços na aprendizagem dos agentes. A Figura 1 apresenta o comportamento do agente durante a fase de avaliação em uma das execuções, dividida em diferentes segmentos (Collectables até Collectables P7).

No final da execução, após aproximadamente 5.500 episódios de teste, o agente alcançou uma média de cerca de 40 minerais coletados. O gráfico está suavizado com uma média móvel de 50 episódios, permitindo observar a tendência de crescimento do desempenho ao longo do tempo, apesar da variabilidade natural do processo de treinamento.

É importante destacar que este resultado ainda é inferior ao desempenho obtido pela DeepMind em seu artigo original (Vinyals et al., 2017), no qual os agentes alcançaram valores significativamente maiores no mesmo ambiente. Essa diferença ressalta tanto a complexidade da tarefa quanto a necessidade de ajustes adicionais na configuração dos algoritmos e nos parâmetros de treinamento. Atualmente, estamos trabalhando em reproduzir a mesma configuração de estado utilizada pela DeepMind em seus experimentos, com o objetivo de verificar se essa abordagem pode melhorar o desempenho do agente em nossos testes.

Resultados Preliminares – *BuildMarines* e *DefeatRoaches*

No cenário *BuildMarines*, os resultados iniciais mostraram uma grande diferença entre o desempenho do agente durante o treino e durante o processo de avaliação. Enquanto no treino os valores de recompensa se mantinham relativamente altos e estáveis, no processo de avaliação o agente não apresentava o mesmo nível de desempenho. Essa disparidade também foi observada no minigame *DefeatRoaches*.

Como primeira tentativa de reduzir esse problema, aplicamos a técnica de mascaramento de ações. No caso do *BuildMarines*, essa modificação trouxe uma melhora visível no desempenho do agente, embora a diferença entre treino e avaliação ainda permaneça, conforme ilustrado nas Figuras 2 e 3.

Uma possível explicação para essa disparidade está relacionada à forma como a recompensa é definida nesses ambientes. No caso do *BuildMarines*, por exemplo, a ação de construir um *marine* só resulta em recompensa muito tempo depois de ter sido executada. Isso gera um forte *delay* na recompensa, dificultando o processo de aprendizagem.

Uma alternativa que estamos considerando é atribuir a recompensa apenas ao final do episódio, tornando-a mais esparsa. Embora essa abordagem aumente consideravelmente o tempo de treinamento necessário, ela pode reduzir o ruído durante o aprendizado e melhorar a coerência entre treino e avaliação.

Discussão Geral

Os resultados alcançados até aqui demonstram que a URNAI 2.0 está cada vez mais apta a servir como base para experimentos robustos em Aprendizado por Reforço Profundo. A modularidade da ferramenta, aliada à integração com bibliotecas consolidadas e ao uso de boas práticas de engenharia de software, garante que a plataforma seja escalável, flexível e confiável.

Embora a reprodução completa dos experimentos da DeepMind ainda não tenha sido concluída, os avanços obtidos indicam que a URNAI está no caminho certo. A possibilidade de aplicar algoritmos da SB3 diretamente sobre ambientes complexos como *StarCraft II* representa uma conquista importante, pois aproxima a plataforma de padrões internacionais de pesquisa e amplia sua relevância para a comunidade científica.

Conclusão

O desenvolvimento da URNAI 2.0 ao longo deste período representou avanços significativos na consolidação da plataforma como uma ferramenta robusta e flexível para experimentação em Aprendizado por Reforço Profundo. A integração prática com a *Stable Baselines3* (SB3) ampliou as possibilidades de uso da ferramenta, permitindo que algoritmos já consolidados na comunidade pudessem ser aplicados de forma direta. Além disso, a adaptação da plataforma para o suporte a *Action Masking* reforçou sua compatibilidade com modelos modernos da SB3, oferecendo maior eficiência no processo de aprendizado.

Outro ponto relevante foi a adoção do *Docker*, que trouxe portabilidade e reprodutibilidade, assegurando que o software pudesse ser executado de maneira consistente em diferentes ambientes. As práticas de engenharia de software — como versionamento estruturado no *GitHub*, uso de *linters* e testes unitários com alta cobertura — foram fundamentais para manter a qualidade do código e garantir que cada componente evoluísse de forma confiável.

No campo experimental, o trabalho encontra-se em progresso na reprodução dos resultados obtidos pela DeepMind em minigames de *StarCraft II*. Embora ainda não tenha sido alcançada a replicação completa, os avanços obtidos até aqui validam a integração da plataforma com o SB3, o registro automático de métricas via *wandb* e a utilização de *Action Masking*, indicando que a ferramenta já está pronta para conduzir experimentos cada vez mais complexos.

Em síntese, a URNAI 2.0 está cada vez mais preparada para apoiar pesquisas em ambientes de alta complexidade, oferecendo uma base sólida que alia qualidade de software, modularidade e integração com ferramentas de ponta. Os próximos passos envolvem não apenas a conclusão da reprodução dos experimentos da DeepMind, mas também a expansão do uso da plataforma em diferentes cenários e a exploração de novas estratégias de aprendizado por reforço profundo.

Referências

- Gibney, E. (2016). Google AI algorithm masters ancient game of Go. *Nature News*, 529(7587), 445.
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep Learning*. MIT Press.
- Silver, D., Schrittwieser, J., Simonyan, K., et al. (2016). Mastering the game of Go with deep neural networks and tree search. *Nature*, 529(7587), 484-489.
- Vinyals, O., Babuschkin, I., Czarnecki, W. M., et al. (2017). *StarCraft II: A New Challenge for Reinforcement Learning*. arXiv preprint arXiv:1708.04782.
- Madeira, C. A. G., et al. (2020). URNAI Tools: A Modular Platform for Deep Reinforcement Learning in Complex Environments. In *Proceedings of the 19th Brazilian Symposium on Computer Games and Digital Entertainment (SBGames 2020)*.
- Thiollent, M. (2011). *Metodologia da pesquisa-ação*. São Paulo: Cortez.

Anexos

Figura 1 - Resultados preliminares do desempenho do agente no minigame CollectMineralShards, mostrando a evolução da média de minerais coletados ao longo dos episódios de avaliação (com média móvel de 50 episódios).

Figura 2 - Recompensa total no processo de avaliação do agente no cenário BuildMarines.

Figura 3 - Recompensa total no processo de treino do agente no cenário BuildMarines.

CÓDIGO: ET1574

AUTOR: GUSTAVO GUILHERME WANDERLEY

COAUTOR: DANIEL LOPES MARTINS

COAUTOR: WELLINGTON SILVA DE OLIVEIRA

COAUTOR: DIEGO MEDEIROS PONTE

ORIENTADOR: ADRIAO DUARTE DORIA NETO

TÍTULO: Previsão do nível de um tanque em tempo real utilizando IA embarcada em um ESP32-CAM-S3 através da plataforma Edge Impulse e TensorFlow lite

Resumo

Este artigo apresenta a implementação de um modelo de visão computacional embarcado no microcontrolador ESP32-CAM para medir o nível de um tanque. O sistema utiliza a plataforma Edge Impulse para o treinamento do modelo e TensorFlow Lite para sua execução. O modelo, baseado em redes neurais convolucionais, foi quantizado para garantir eficiência em dispositivos com recursos limitados. O desempenho do sistema de IA é rigorosamente avaliado por meio de uma comparação direta com um sensor de nível industrial de alta precisão, validando sua precisão e identificando suas limitações em um cenário prático.

Palavras-chave: IA, GPU, Quantização, ESP32, Visão Computacional, Edge Impulse, Tensor

TITLE: Real-time tank level prediction using AI embedded in an ESP32-CAM-S3 through the Edge Impulse platform and TensorFlow Lite

Abstract

This paper presents the implementation of an embedded computer vision model on the ESP32-CAM microcontroller to measure a tank's level. The system uses the Edge Impulse platform for model training and TensorFlow Lite for its execution. The model, based on convolutional neural networks, was quantized to ensure efficiency on resource-constrained devices. The AI system's performance is rigorously evaluated through a direct comparison with a high-precision industrial level sensor, validating its accuracy and identifying its limitations in a practical scenario.

Keywords: AI, GPU, Quantization, ESP32, Computer Vision, Edge Impulse, Tensor

Introdução

A medição de nível de tanques é um pilar para a automação industrial, onde sensores de alta precisão, como os de pressão e radar, são essenciais para garantir segurança e eficiência. Essas tecnologias são a base para o controle de processos críticos, oferecendo alta confiabilidade na maioria das aplicações.

Contudo, existem cenários operacionais específicos onde até mesmo esses sensores robustos enfrentam desafios. Um sensor de pressão, por ser inerentemente invasivo, pode sofrer com entupimentos ou se degradar rapidamente com produtos químicos, tornando-se um ponto de falha \cite{b2}. Já os sensores de radar, apesar de sua excelência, podem ser temporariamente "cegados" por uma camada de espuma muito densa ou por um material que absorva as ondas eletromagnéticas, limitações conhecidas na literatura \cite{b9}.

É para preencher essas lacunas que a medição por câmera e Inteligência Artificial se apresenta como uma solução estratégica e complementar. O sistema utiliza uma câmera de baixo custo para "ler" um indicador de nível externo, funcionando como um par de olhos digital, independente e verificador. Por não ter contato com o produto, é imune a entupimentos e corrosão.

O verdadeiro impacto desta tecnologia está em sua capacidade de trabalhar em conjunto com os sistemas existentes. Ela pode atuar como um sistema de redundância de baixo custo, validando as leituras do sensor principal e alertando sobre divergências que possam indicar uma falha iminente. Adicionalmente, permite expandir a automação para tanques auxiliares de forma econômica, garantindo uma visão completa e mais segura de toda a planta industrial \cite{b2}.

Com o avanço da Indústria 4.0 e da Internet das Coisas (IoT), surge uma crescente demanda por soluções de monitoramento que sejam mais flexíveis, de baixo custo e não invasivas. Neste contexto, a aplicação de Inteligência Artificial (IA) embarcada em microcontroladores de baixo custo apresenta-se como uma alternativa promissora. A capacidade de inferir uma variável física, como o nível, a partir de imagens capturadas em tempo real, elimina a necessidade de contato direto com o fluido, sendo viabilizada por frameworks como o TensorFlow Lite, projetado especificamente para essa classe de dispositivos \cite{b1}.

A relevância deste tipo de pesquisa é acentuada pela \mbox{atual} necessidade de três tendências tecnológicas. Primeiro, a presença de microcontroladores potentes e acessíveis, como a série ESP32, que integram processamento, conectividade e interfaces de câmera a um custo acessível. Segundo, a maturidade de frameworks de software, como o TensorFlow Lite, que são bem otimizados para hardware com recursos restritos. Terceiro, o surgimento de plataformas, como a Edge Impulse, que facilitam o uso e aceleram o ciclo de desenvolvimento da IA embarcada. É a união destes fatores que torna a presente proposta não apenas viável, mas também torna possível solucionar muitos problemas com esta tecnologia.

Apesar desses pontos positivos, a validação rigorosa de tais sistemas de baixo custo contra padrões industriais consolidados é um passo crucial para garantir sua confiabilidade no chão de fábrica. Portanto, o objetivo central deste trabalho não é apenas demonstrar a viabilidade técnica da solução, mas, fundamentalmente, avaliar seu desempenho de forma quantitativa, comparando diretamente as medições da IA com as de um sensor de nível industrial de alta precisão.

Metodologia

Fundamentação Teórica

Para contextualizar a solução desenvolvida, esta seção aborda as metodologias de medição, os conceitos de software e o hardware que viabilizam a implementação da inteligência artificial.

Metodologias de Medição de Nível

Abordagem Industrial com Sensores Dedicados A medição de nível em ambientes industriais é classicamente realizada por meio de instrumentação dedicada, projetada para alta precisão e robustez. Exemplos comuns incluem sensores ultrassônicos que funcionam emitindo um pulso sonoro e medindo o tempo que esse pulso leva para refletir de volta após atingir a superfície desejada. A distância é calculada com base nesse tempo de voo, considerando a velocidade do som no meio em questão. Outro exemplo são os sensores de radar que operam de maneira semelhante, mas utilizam ondas eletromagnéticas em vez de som. Por conta disso, eles são menos suscetíveis a interferências ambientais e oferecem maior precisão e alcance, além de funcionarem melhor em ambientes com vapor, poeira ou pressões extremas. Outra abordagem comum, e a utilizada como referência neste trabalho, emprega sensores de pressão estática (manométrica) instalados na base do tanque. Estes sensores medem a pressão exercida pela coluna de líquido para então inferir a altura. O sistema de referência ilustrado na Figura, representa essa abordagem convencional, utilizando sensores industriais (Sensor Nível 1 e 2) para obter leituras diretas e precisas.

Planta do sistema de recalque, ilustrando o uso de sensores de nível industriais dedicados.

Abordagem Proposta com Visão Computacional Como alternativa, este estudo propõe uma abordagem não invasiva e de baixo custo que utiliza visão computacional para inferir o nível do tanque. A estrutura experimental para esta metodologia é esquematizada na Figura. Um microcontrolador com câmera, o ESP32-CAM-S3, é posicionado externamente ao tanque, utilizando seu campo de visão para capturar imagens de um indicador visual de nível. Essas imagens são pré-processadas localmente pelo próprio hardware do ESP32-CAM-S3, de forma a adequá-las aos parâmetros exigidos pelo modelo. Em seguida, um modelo de inteligência artificial embarcado realiza a inferência do nível do tanque com base nas imagens processadas, eliminando a necessidade de sensores físicos dedicados à medição de nível.

Esquemático da estrutura proposta, com o ESP32-CAM-S3 medindo o nível do tanque através de seu campo de visão.

Visão Computacional e Redes Neurais Convolucionais A Visão Computacional é um campo da IA que busca capacitar os computadores a enxergar e interpretar o mundo visual. As tarefas incluem classificação de imagens, detecção de objetos e, como neste trabalho, regressão para estimar um valor a partir de uma imagem. As Redes Neurais Convolucionais (CNNs) são a arquitetura de deep learning predominante para essas tarefas.

Uma CNN é tipicamente composta por três tipos de camadas: convolucionais, de pooling e totalmente conectadas.

A camada convolucional é o núcleo de uma CNN. Nela, utilizamos filtros (kernels) que são pequenas matrizes com valores numéricos. Esses filtros percorrem a imagem de entrada, posição por posição, realizando uma operação chamada produto escalar (ou produto de ponto) entre os valores do filtro e os valores da imagem naquela região. Esse processo gera um mapa de características (feature map) que destaca padrões específicos na imagem, como

arestas, texturas ou formas. As primeiras camadas aprendem características simples, como bordas e texturas. À medida que a rede se aprofunda, as camadas seguintes, geralmente compostas por um número maior de filtros, combinam essas informações para identificar padrões mais complexos. Isso ocorre porque uma CNN pode conter várias camadas convolucionais empilhadas, e cada camada extrai representações cada vez mais abstratas a partir das ativações da anterior. Dessa forma, a rede é capaz de construir gradualmente uma compreensão mais rica da imagem, indo de elementos básicos até estruturas mais complexas.

Inserida entre as camadas convolucionais, a camada de pooling (geralmente Max Pooling) tem o objetivo de reduzir a dimensão dos mapas de características. Ela opera sobre pequenas janelas do mapa e substitui os valores daquela janela por um único valor (o máximo, no caso do Max Pooling). Isso torna a representação mais robusta a pequenas translações na imagem e reduz a carga computacional da rede.

Após as sucessivas etapas de convolução e pooling, os dados são processados por uma camada de achatamento, conhecida como Flatten. A função desta camada é transformar a estrutura de dados multidimensional (por exemplo, uma matriz de características de 2D) em um vetor unidimensional. Essa reorganização é um passo fundamental para que os dados possam ser utilizados pela rede, composta por uma ou mais camadas totalmente conectadas (Fully Connected ou Dense), onde cada neurônio se conecta a todos os neurônios da camada anterior.

Estas camadas densas finais funcionam como a parte decisória da rede, atuando como um classificador ou regressor que utiliza as características extraídas para realizar a predição. A configuração da última camada densa, especialmente sua função de ativação, define a natureza da saída do modelo:

Regressão (Saída Linear): Para prever um valor contínuo, a última camada possui uma saída linear. Isso permite que a rede produza qualquer valor numérico. É o caso deste trabalho, onde o objetivo é a regressão do nível do tanque, prevendo um valor específico dentro de um intervalo contínuo.

Classificação Binária (Função Sigmoid): Para problemas de classificação binária (onde a resposta é "sim" ou "não"), utiliza-se a função de ativação Sigmoid. Ela comprime a saída para um valor entre 0 e 1, que pode ser interpretado como a probabilidade de pertencer a uma classe.

Resultados e Discussões

Classificação Multiclasse (Função Softmax): Para classificar uma entrada em três ou mais categorias, a função Softmax é empregada. Ela converte as saídas da camada em um vetor de probabilidades, onde a soma de todos os elementos é 1. Cada valor no vetor representa a probabilidade de a entrada pertencer a uma das classes. Por exemplo, um modelo que classifica imagens de veículos entre "carro", "moto" e "caminhão".

A arquitetura específica do modelo utilizado neste trabalho, ilustrada na Figura, foi desenvolvida com o auxílio da plataforma Edge Impulse. Uma característica notável e não convencional desta arquitetura é a utilização de um kernel de dimensões 128x128 na primeira camada convolucional. É crucial destacar que esta escolha diverge significativamente das práticas habituais em Redes Neurais Convolucionais (CNNs), que faz o uso de kernels

pequenos (como 3x3 ou 5x5) e empilhados em várias camadas. A justificativa para a adoção deste kernel de grande dimensão reside no fato de que esta arquitetura é a forma padrão da ferramenta, a qual correspondeu bem em testes estáticos, somado à falta de tempo para testar outros tamanhos de kernels até a data de submissão do trabalho. Essa escolha, embora atípica, tem uma consequência direta: um kernel de 128x128 resulta em um número exponencialmente maior de parâmetros a serem treinados já na primeira camada, o que contribui significativamente para o tamanho final do modelo. Apesar da divergência, é importante ressaltar que o modelo gerado, mesmo com essa particularidade, conseguiu realizar a predição do nível do tanque.

Arquitetura do modelo de CNN utilizado.

O Framework TensorFlow Lite e a Otimização de Modelos Historicamente, a execução de modelos de deep learning era restrita a servidores com GPUs potentes. O TensorFlow Lite (TFLite) é o framework do Google projetado para quebrar essa barreira, permitindo que modelos de IA rodem em dispositivos com recursos restritos. Sua variante, TensorFlow Lite for Microcontrollers, é otimizada para ambientes com memória na ordem de kilobytes.

O uso de TFLite para processamento local (on-device) oferece benefícios importantes como maior privacidade, baixa latência, baixo consumo de energia e custo reduzido de hardware.

Para viabilizar a execução, os modelos precisam ser otimizados. A técnica central é a quantização, que consiste em reduzir a precisão numérica dos pesos e/ou ativações da rede. Neste trabalho, foi convertido os pesos de float32 para int8. Essa otimização reduz o tamanho do modelo em aproximadamente 4x e acelera a inferência, o que poderá ser melhor visto na seção de metodologia experimental.

Hardware da Solução Proposta Para a implementação da metodologia de visão computacional (Figura), o hardware selecionado é um componente central. O módulo utilizado foi o ESP32-CAM-S3, uma placa de desenvolvimento que utiliza o poderoso microcontrolador ESP32-S3 da Espressif Systems.

O ESP32 é um microcontrolador de baixo custo e alta performance, conhecido por sua versatilidade em projetos de Internet das Coisas (IoT). Em sua essência, ele integra um processador, memória e uma vasta gama de periféricos em um único chip. Isso inclui:

Pinos de I/O(GPIO): Permitem a comunicação com sensores, atuadores e outros componentes eletrônicos.

Conversores Analógico-Digital (ADC) e Digital-Analógico (DAC): Essenciais para a interface com sensores que fornecem sinais analógicos.

Sensores de Toque Capacitivo: Possibilitam a criação de interfaces de usuário sensíveis ao toque.

Interfaces de Comunicação: Suporte nativo para Wi-Fi e Bluetooth, além de protocolos como SPI, I2C e UART.

A versão S3 do ESP32, utilizada neste projeto é uma evolução projetada especificamente para tarefas mais exigentes, como o uso de modelos de Inteligência Artificial. Suas principais vantagens são:

Processador Dual-Core Xtensa® LX7: Mais moderno e eficiente que o LX6 das gerações anteriores.

Extensões Vetoriais (Vector Instructions): Este é o grande diferencial. O processador LX7 inclui instruções específicas para acelerar operações matemáticas em vetores e matrizes. Como as redes neurais, base do TensorFlow Lite, dependem massivamente de cálculos matriciais (multiplicações e convoluções), essas extensões aceleram drasticamente o tempo de inferência do modelo.

Boa quantidade de memória: Com 8 MB de PSRAM e 16 MB de Flash, o ESP32-CAM-S3 possui os recursos necessários para armazenar e executar modelos de IA mais complexos, que não caberiam em microcontroladores mais simples.

A combinação de um processador mais rápido com aceleração de hardware específica para IA torna o ESP32-S3 a escolha ideal para executar modelos otimizados com o framework TensorFlow Lite.

A tabela a seguir resume as especificações do módulo empregado.

Especificações do Hardware ESP32-CAM S3 Especificação: Microcontrolador Descrição: ESP32-S3, Dual-core Xtensa® 32-bit LX7

Especificação: Memória Flash Descrição: 16 MB (SPI Flash)

Especificação: PSRAM Descrição: 8 MB

Especificação: Câmera Descrição: OV2640 / OV5640 suportadas

Especificação: Conectividade Descrição: Wi-Fi 802.11 b/g/n + Bluetooth 5.0 (LE)

Metodologia Experimental

Construção do Dataset A base para o modelo de IA foi um conjunto de 1240 imagens no formato RGB (320x240 pixels), capturadas com a própria câmera do ESP32-CAM-S3. Para garantir diversidade, foram coletadas 40 fotos para cada centímetro de nível do tanque, com objetivo de representar variações naturais que ocorrem durante a captura, como mudanças na iluminação ambiente, presença de ruído na imagem e pequenas variações no foco. Essas variações ajudam a tornar o modelo mais robusto, evitando que ele se ajuste apenas a condições ideais e permitindo uma melhor generalização durante a inferência. O conjunto foi dividido de forma aleatória, com 75% das imagens para treinamento e 25% para teste.

Treinamento e Quantização O treinamento foi conduzido na plataforma Edge Impulse. O modelo de CNN (Figura) foi treinado por 150 épocas utilizando o otimizador Adam com uma taxa de aprendizado de 0.005 e um tamanho de lote (batch size) de 128. Essa parametrização foi feita com base em uma série de testes. As métricas de erro como, MSE (Erro Quadrático Médio) e o MAE (Erro absoluto Médio) foram essenciais para encontrar esta combinação. Uma porção dos dados de treino foi usada como conjunto de validação para monitorar o desempenho e evitar sobreajuste (overfitting). Após o treinamento, o modelo foi convertido e quantizado para o formato TensorFlow Lite (TFLite) com pesos em int8. Inicialmente foi utilizado o ESP32-S para tirar resultados iniciais e posteriormente, com a aquisição do ESP32-CAM-S3, o trabalho passou a utilizá-lo pelo ganho de performance que será apresentado nos tópicos seguintes do artigo. A seguir temos a Tabela que representa as

especificações do ESP32-S e a Tabela representa a diferença de desempenho entre o modelo quantizado(int8) e o modelo sem otimização(float32).

Especificações do Hardware ESP32-S Especificação: Microcontrolador Descrição: ESP32-S, Dual-core Tensilica Xtensa® LX6

Especificação: Frequência de Clock Descrição: Até 240 MHz

Especificação: Memória Flash Descrição: Tipicamente 4 MB (SPI Flash no módulo)

Especificação: SRAM Descrição: 520 KB (Integrada ao chip)

Especificação: Conectividade Descrição: Wi-Fi 802.11 b/g/n + Bluetooth v4.2

Comparação entre modelo quantizado e float32(ESP32-S) Parâmetro: Latência de Inferência Quantizado (int8): 1.546 ms Float32: 8.584 ms

Parâmetro: Uso de RAM Quantizado (int8): 322.6 KB Float32: 1.3 MB

Parâmetro: Uso de Flash (ROM) Quantizado (int8): 63.2 KB Float32: 168.4 KB

Parâmetro: Acurácia (teste) Quantizado (int8): 99.2% Float32: 99.8%

Como apresentado na Tabela, temos a diferença de desempenho entre os modelos numa simulação baseada no ESP32-S. Como já mencionado previamente, o resultado obtido com o chip mais potente será apresentado ao longo do artigo.

Resultados e Discussão Esta seção detalha os resultados obtidos, da análise do modelo à validação prática contra o sensor industrial.

Análise da Extração de Características

A Figura apresenta uma visualização UMAP (Uniform Manifold Approximation and Projection) das características extraídas pela CNN. A utilização do UMAP justifica-se por ser uma técnica de redução de dimensionalidade não linear, especialmente eficaz em preservar a estrutura topológica de dados complexos. Ao projetar as representações internas de alta dimensão aprendidas pela rede para um espaço bidimensional, o UMAP permite uma inspeção visual precisa da geometria do conhecimento adquirido pelo modelo.

O gráfico resultante fornece evidências claras sobre o sucesso da aprendizagem. Observa-se a formação de agrupamentos distintos, indicando que a CNN aprendeu um espaço de características onde amostras de diferentes níveis de água são altamente separáveis. Além disso, a progressão lógica que vai do vermelho (nível 0) ao azul (nível 30), demonstra que o modelo não apenas classificou, mas também ordenou as representações de maneira análoga à natureza física do problema.

Visualização UMAP das características extraídas pelo modelo no conjunto de teste. As cores representam os diferentes níveis do tanque, de 0 (vermelho) a 30 (azul).

Desempenho do Modelo em Ambiente de Teste O modelo quantizado foi avaliado no conjunto de teste. Conforme a Tabela, o modelo alcançou um Erro Absoluto Médio (MAE) de apenas 0.45 cm e uma Raiz do Erro Quadrático Médio (RMSE) de 0.57 cm. Estes valores, obtidos em um ambiente sem perdas, demonstram a alta precisão teórica do modelo. A Figura ilustra essa correlação.

Resumo do Desempenho do Modelo no Conjunto de Teste Métrica: Erro Quadrático Médio (MSE) Valor (Teste): 0.32

Métrica: Erro Absoluto Médio (MAE) Valor (Teste): 0.45 cm

Métrica: Raiz do Erro Quadrático Médio (RMSE) Valor (Teste): 0.57 cm

Gráfico de dispersão gerado pelo Edge Impulse dos resultados no conjunto de teste, comparando o nível real com a predição do modelo.

Validação Prática e Análise do Erro O teste definitivo foi a validação em cenário real, comparando as previsões do ESP32-CAM-S3 com as medições do sensor YOKOGAWA EJX110B, conforme a Figura. O Erro Absoluto Médio observado foi de 3.63 cm, significativamente maior que o erro teórico. Essa discrepância evidencia um problema em que as condições do mundo real (iluminação, posicionamento da câmera) diferem das do dataset de treinamento. Uma análise mais atenta do erro na Figura revela um comportamento não linear. Para níveis baixos (9-12 cm), o erro é positivo e relativamente baixo, indicando uma ligeira superestimação. Contudo, para níveis altos (20-27 cm), o erro torna-se significativamente negativo, demonstrando uma forte subestimação do modelo. Isso pode sugerir que as características visuais em níveis mais altos, talvez devido a reflexos ou distorções da lente, não foram bem representadas no dataset de treinamento, ou que o modelo teve dificuldade em generalizar para essa faixa. Outro ponto a ser destacado é que a quantização leva a perdas na inferência do modelo quando este é inserido no sistema microcontrolado, que neste caso, é o ESP32-CAM-S3.

Gráfico comparativo entre as medições do ESP32-CAM-S3 e do sensor industrial.

Redução do tempo de inferência na prática

Como apresentado na Tabela, o modelo teve um desempenho bem melhor do que foi apresentado na Tabela justamente pela diferença significativa de desempenho entre o chip ESP32-S e o ESP32-S3. A PSRAM e a memória Flash não demonstraram ser o gargalo que causou a limitação de desempenho observada. Contudo, para modelos maiores e de maior complexidade, o ESP32-S3 [7] oferece recursos específicos que o tornam mais eficiente para a execução de redes neurais embarcadas. Entre esses recursos, destaca-se o suporte a operações de ponto flutuante que contribui para uma maior precisão nos cálculos realizados durante a inferência dos modelos. Essas características permitem que o ESP32-S3 realize o processamento local de modelos de inteligência artificial com maior eficiência, reduzindo a latência e a dependência de comunicação externa para o processamento dos dados. Dessa forma, ele se mostra adequado para aplicações que exigem execução rápida e precisa de redes neurais em ambientes embarcados.

Desempenho modelo quantizado (ESPS3-CAM) Parâmetro: Latência de Inferência Quantizado (int8): 283.28 ms

Parâmetro: Uso de RAM Quantizado (int8): 322.6 KB

Parâmetro: Uso de Flash (ROM) Quantizado (int8): 63.2 KB

Parâmetro: Acurácia (teste) Quantizado (int8): 99.2%

A latência de inferência de aproximadamente 283 ms apresentada na Tabela demonstra que o modelo quantizado possui tempo de resposta suficientemente rápido para aplicações em sistemas de controle embarcado. O trabalho apresenta uma alternativa viável frente aos sensores industriais tradicionais, principalmente quando se considera o custo significativamente reduzido, a velocidade de inferência compatível com requisitos de tempo real e a possibilidade de uso redundante, agregando segurança ao sistema. Tal abordagem possibilita a aplicabilidade no mundo real e abre espaço para adoção em cenários onde a relação custo-benefício é decisiva.

Trabalhos Futuros Os resultados sugerem diversas vias para aprimoramento, focadas tanto em acurácia quanto em eficiência.

Melhorias de Acurácia e Robustez

Aumento e Realimentação do Dataset: Coletar mais imagens sob as diversas condições de operação.

Diminuição do Kernel

O trabalho utilizou o kernel padrão de 128x128 gerado pela plataforma Edge Impulse. Embora funcional, essa configuração não representa necessariamente a solução mais otimizada. Devido à falta de tempo e aos bons resultados obtidos nos testes estáticos, não foram testados outros tamanhos mais convencionais (como 3x3 ou 5x5). No entanto, para trabalhos futuros, espera-se que uma abordagem convencional junto do aumento de camadas densas, possa melhorar o desempenho e a eficiência do modelo.

Otimização de Eficiência e Desempenho

Aceleração por Hardware: Otimizar o código para usar com mais eficiência as instruções de aceleração de IA específicas do ESP32-S3.

Conclusão

Conclusão Este trabalho demonstrou a viabilidade de usar um modelo de visão computacional em um microcontrolador ESP32-CAM-S3 para medir o nível de um tanque em tempo real. A aplicação de técnicas de otimização com o framework TensorFlow Lite, como a quantização, foi crucial para criar um sistema com baixa latência (283.28 ms) e consumo de memória (322.6 KB), que alcançou um erro teórico de apenas 0.45 cm.

O ponto-chave do estudo, no entanto, foi a possibilidade de aplicação prática. A comparação com um sensor industrial YOKOGAWA revelou um erro médio de 3,63 cm e um desvio não linear. Essa diferença não invalida a abordagem, pois há bastante margem para melhorias e também ajuda a definir seu melhor contexto de aplicação: para cenários onde uma margem de erro de centímetros é aceitável, ou onde a redundância é mais importante que a precisão absoluta, a solução com IA é uma alternativa de custo muito inferior, não invasiva e de rápida implementação.

Em resumo, o projeto mostra um uso promissor do TensorFlow Lite em microcontroladores como uma ferramenta poderosa para a democratização da automação industrial. Seu impacto vai além da simples substituição de sensores. Ele permite a criação de sistemas de

monitoramento redundantes que aumentam a segurança. A tecnologia representa, assim, um passo importante para a criação de plantas industriais mais inteligentes, flexíveis e resilientes.

Agradecimentos

A presente pesquisa foi realizada com o apoio do Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq) e da Universidade Federal do Rio Grande do Norte (UFRN), por meio da concessão de bolsa de iniciação científica (IC) durante o desenvolvimento deste projeto.

Referências

- R. David, J. Duke, A. Jain, V. J. Reddi, N. Shpancer, D. S. G. Sampedro, et al., "TensorFlow Lite Micro: Embedded Machine Learning on TinyML Systems," in Proc. MLSys Conf., 2021.
- B. G. Lipták, Ed., Instrument Engineers' Handbook, Vol. 1: Process Measurement and Analysis, 4th ed. Boca Raton, FL, USA: CRC Press, 2003.
- R. Szeliski, Computer Vision: Algorithms and Applications, 2nd ed. Cham, Switzerland: Springer, 2022.
- I. Goodfellow, Y. Bengio, and A. Courville, Deep Learning. Cambridge, MA, USA: MIT Press, 2016.
- P. A. R. V. B. da Silva, A. G. S. Filho, and E. J. P. da Silva, "Water level measurement in an open channel using a computer vision system," in Proc. IEEE Int. Instrum. Meas. Technol. Conf. (I2MTC), 2019, pp. 1-6.
- B. Jacob et al., "Quantization and Training of Neural Networks for Efficient Integer-Arithmetic-Only Inference," in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), 2018, pp. 2704-2713.
- Espressif Systems, ESP32-S3 Series Datasheet, Ver. 2.0, Mar. 2025. [Online]. Available: https://www.espressif.com/sites/default/files/documentation/esp32-s3_datasheet_en.pdf. [Accessed: May 28, 2025].
- Ai-Thinker, ESP-32S Product Specification, Aug. 2025. [Online]. Available: <https://docs.ai-thinker.com/en/esp32/spec/esp32s>. [Accessed: May 28, 2025].
- J. M. M. de Freitas, "Sensor de nível por micro-ondas e tecnologia RADAR-FMCW," Dissertação de Mestrado, Univ. Federal de Itajubá (UNIFEI), Itajubá, Brasil, 2013. [Online]. Available: <https://repositorio.unifei.edu.br/jspui/handle/123456789/922>
- L. McInnes, J. Healy, and J. Melville, "UMAP: Uniform Manifold Approximation and Projection for Dimension Reduction," Journal of Open Source Software, vol. 3, no. 29, p. 861, 2018. [Online]. Available: <https://joss.theoj.org/papers/10.21105/joss.00861>

Anexos

plot de resultados 2

arquitetura do modelo

gráfico da função UMAP

plot dos resultados

diagrama de sensores industriais

CÓDIGO: ET1616

AUTOR: DAVI WHAGNER DE LIMA HOFFMANN

COAUTOR: STEFANE DE ASSIS ORICHUELA

ORIENTADOR: ALYSON MATHEUS DE CARVALHO SOUZA

TÍTULO: VRPG: Suporte de auxílio imersivo acessível e de baixo custo para campanhas de RPG de mesa.

Resumo

Os jogos de interpretação de papéis de mesa (RPG) constituem atividades colaborativas que favorecem a criatividade, a interação social e a imersão narrativa. A Realidade Virtual (RV) surge como um recurso promissor para ampliar a sensação de presença espacial nessas campanhas. No entanto, as soluções disponíveis ainda são limitadas pelo alto custo de hardware e pelo foco em experiências remotas. Este trabalho de pesquisa apresenta o VRPG, um sistema imersivo de baixo custo que integra RV em campanhas presenciais de RPG por meio do uso de smartphones e visores Google Cardboard. O sistema valoriza o reaproveitamento de equipamentos já disponíveis para jogos, promovendo pesquisas sobre maior acessibilidade e potencial de utilização também como ferramenta educacional. No VRPG, o mestre conduz os avatares dos jogadores em um mapa 3D compartilhado, enquanto os participantes exploram o ambiente de forma imersiva. A aplicação foi desenvolvida em Unity e utiliza comunicação peer-to-peer em rede local, garantindo acessibilidade e sincronização em tempo real. Como etapa futura, serão realizados testes empíricos para avaliar o impacto da imersão em RV na experiência, engajamento e criatividade dos jogadores durante as sessões de RPG, bem como seu uso em contextos educativos.

Palavras-chave: Realidade Virtual. RPG. Multiusuário.

TITLE: VRPG: Accessible and Low-Cost Immersive Support for Tabletop RPG Campaigns.

Abstract

Tabletop role-playing games (RPGs) are collaborative activities that foster creativity, social interaction, and narrative immersion. Virtual Reality (VR) emerges as a promising resource to enhance the sense of spatial presence in such campaigns. However, current solutions remain limited due to the high cost of hardware and the focus on remote experiences. This research presents VRPG, a low-cost immersive system that integrates VR into face-to-face RPG campaigns through the use of smartphones and Google Cardboard headsets. The system emphasizes the reuse of equipment already available for gaming, promoting research on greater accessibility and potential use as an educational tool. In VRPG, the game master controls the players' avatars on a shared 3D map, while participants explore the environment

immersively. The application was developed in Unity and uses peer-to-peer communication over a local network, ensuring accessibility and real-time synchronization. As a future step, empirical tests will be conducted to evaluate the impact of VR immersion on players' experience, engagement, and creativity during RPG sessions, as well as its application in educational contexts.

Keywords: Virtual Reality. RPG. Multiuser.

Introdução

Os jogos de interpretação de papéis (Role-Playing Games – RPG) de mesa são atividades narrativas colaborativas nas quais os participantes assumem personagens fictícios e interagem sob a mediação de um mestre, que controla os eventos e cenários da história. As campanhas de RPG desempenham papel relevante na promoção de habilidades como criatividade, pensamento estratégico e colaboração social, além de possibilitarem narrativas prolongadas que podem se estender por semanas ou meses.

O avanço das tecnologias imersivas, especialmente da Realidade Virtual (RV), tem ampliado as formas de interação em experiências narrativas e educacionais. Ambientes tridimensionais virtuais permitem maior percepção espacial e sensação de presença, beneficiando a compreensão de contextos e a imersão do usuário. No entanto, a aplicação de RV em campanhas de RPG de mesa ainda é restrita por fatores como custo de hardware e complexidade técnica, dificultando sua adoção por grupos que realizam partidas presenciais.

Pesquisas anteriores demonstram o potencial da RV para enriquecer o RPG. Estudos com dispositivos de alto desempenho, como o Meta Quest 2, permitem a criação de avatares e interações remotas entre jogadores e mestres em ambientes virtuais compartilhados. Outras investigações apontam que elementos sonoros desempenham papel crucial na ampliação da sensação de presença e no engajamento narrativo durante as sessões. Apesar desses avanços, permanece a necessidade de soluções mais acessíveis e voltadas para campanhas presenciais.

Este trabalho apresenta o VRPG, um sistema de suporte imersivo de baixo custo para campanhas de RPG de mesa. A solução utiliza smartphones acoplados a visores de realidade virtual baseados em Google Cardboard e é projetada para partidas locais, com comunicação entre jogadores e mestre via rede local. O sistema permite que o mestre controle o posicionamento dos jogadores em um mapa tridimensional enquanto estes visualizam o ambiente de forma imersiva. O objetivo é oferecer uma alternativa viável e economicamente acessível para grupos interessados em integrar elementos de RV às suas campanhas presenciais.

Metodologia

A aplicação foi concebida como um suporte para mestres e jogadores em campanhas presenciais de RPG, priorizando recursos de baixo custo a fim de ampliar a acessibilidade e facilitar o uso de tecnologias digitais durante as sessões. Para isso, optou-se pelo uso do Google Cardboard, solução acessível e de fácil implementação, em conjunto com smartphones comuns. O desenvolvimento foi realizado no motor de jogos Unity (versão

2022.3.20), adotando um modelo 3D de masmorra disponibilizado na Asset Store como ambientação inicial da campanha.

Embora o ambiente tenha sido inicialmente idealizado para ser gerado de forma procedural e contínua, de acordo com as preferências do mestre e com apoio de inteligência artificial generativa, nesta primeira versão decidiu-se priorizar a implementação da integração multiusuário.

A aplicação foi organizada em duas cenas principais: Mestre e Jogadores.

Na cena do Mestre, a câmera é posicionada acima da masmorra sem o teto, proporcionando uma visão superior que funciona como mapa de controle da campanha.

Na cena dos Jogadores, a câmera é vinculada ao avatar do personagem e configurada para renderização em Realidade Virtual por meio do Google Cardboard.

Ambas as cenas utilizam o mesmo modelo 3D, alinhado nos mesmos eixos de localização para garantir consistência espacial. A única diferença estrutural é a ausência do teto na cena do Mestre, permitindo a visualização completa do interior da masmorra.

Além dos ambientes imersivos, a aplicação conta com um menu inicial (Fig. 1), que orienta os usuários e possibilita tanto o preenchimento da ficha do jogador quanto a criação da sala da campanha. O plano de fundo desse menu foi gerado por inteligência artificial generativa, escolhido para oferecer uma interface visualmente amigável e intuitiva.

A comunicação entre os participantes ocorre exclusivamente em rede local, permitindo que todos os dispositivos conectados compartilhem a experiência sem necessidade de infraestrutura externa. Para isso, foram implementados dois scripts principais: um para o usuário mestre e outro para os usuários jogadores. Quando o mestre inicia uma nova sala, o sistema entra em estado de escuta, aguardando mensagens de status: “NovoJogador” ou “JogadorExistente”. Ao se conectar, um novo jogador envia uma mensagem em formato JSON, que aciona a instanciação de seu avatar no mapa na posição inicial definida pela campanha (Fig. 3). Em seguida, o mestre transmite ao jogador sua posição atual, garantindo a coerência espacial entre os cenários.

Durante a campanha, o mestre pode selecionar e movimentar os avatares no mapa, e essas ações são refletidas em tempo real nos dispositivos dos jogadores, que acompanham a atualização de suas posições dentro da experiência em Realidade Virtual. Essa dinâmica é sustentada por uma arquitetura peer-to-peer, em que os clientes se comunicam diretamente em rede local, com o mestre atuando como nó central de coordenação.

Como perspectiva de evolução, está em desenvolvimento uma estrutura de mensagens estendida, que permitirá ao mestre transmitir periodicamente a lista completa das posições de todos os jogadores a cada participante. Essa funcionalidade tem como objetivo ampliar a sensação de presença coletiva e promover maior interação colaborativa durante as campanhas de RPG.

Resultados e Discussões

Como resultados iniciais, foram aplicados os primeiros testes com alguns participantes, para compreender o quão entendível estava a interface inicial e a experiência do usuário ao entrar na partida.

Um roteador foi utilizado para realizar a transmissão de dados locais entre o mestre e os participantes. O mestre, por sua vez, utilizou um notebook enquanto os dois participantes presentes utilizaram os próprios aparelhos móveis.

Conclusão

Este projeto apresentou o desenvolvimento inicial do VRPG, um software voltado a apoiar a imersão de jogadores em campanhas presenciais de RPG de mesa por meio da Realidade Virtual, destacando-se pela adoção de recursos acessíveis e de baixo custo. Como próximos passos, está prevista a realização de uma primeira sessão experimental com uma campanha de RPG, com o objetivo de avaliar o grau de imersão proporcionado, bem como o impacto na criatividade e na sensação de presença dos participantes, comparando sessões realizadas com e sem o uso da RV.

Referências

COVER, Jennifer Grouling. The creation of narrative in tabletop role-playing games. McFarland, 2014.

PIMENTEL, Daniel; FOXMAN, Maxwell; DAVIS, Donna Z.; MARKOWITZ, David M. Virtually real, but not quite there: Social and economic barriers to meeting virtual reality's true potential for mental health. *Frontiers in Virtual Reality*, v. 2, p. 627059, 2021.

NIARCHOS, Anastasios N.; IOANNIDIS, Yannis; KATIFORI, Akrivi. The use of Immersive VR in Tabletop Role Playing Games as a Virtual Tabletop. 2023.

CÓDIGO: ET1620

AUTOR: LORENZO SANTOS FURTADO

COAUTOR: ANA CATARINA GONÇALVES FUENTES

ORIENTADOR: RENAN CIPRIANO MOIOLI

TÍTULO: Uso de Redes Neurais Pulsadas no Processamento de Imagens de Tumores de Pele

Resumo

Neste trabalho foi desenvolvido e avaliado Redes Neurais Pulsadas (SNNs), um modelo bioinspirado, para a classificação de imagens de câncer de pele. O objetivo do projeto é alcançar a alta precisão das Redes Neurais Convolucionais (CNNs) com um custo computacional e energético significativamente menor. A pesquisa utilizou as bibliotecas pyGeNN e mlGeNN no Google Colab para criar e treinar as SNNs. O desempenho do modelo foi testado com os conjuntos de dados MNIST, DermaMNIST e câncer de pele do ISIC. A rede apresentou uma acurácia de 93,72% no conjunto de teste MNIST e 63% no DermaMNIST. Para as imagens do ISIC foi realizado um pré-processamento minucioso incluindo redimensionamento, ajuste de iluminação e aumento de dados, para otimizar as imagens para a rede. No entanto, o treinamento com o dataset ISC revelou a limitação do ambiente Google Colab, que não suportou a demanda computacional da tarefa. A acurácia máxima obtida durante os treinamentos foi de 65%, o que sublinha o desafio de treinar SNNs em grande escala sem hardware especializado. Os resultados sugerem que, embora as redes pulsadas construídas com o GeNN e suas bibliotecas tenham potencial, sua aplicação em grande escala exige ajustes significativos nos hiperparâmetros e, principalmente, a utilização de uma infraestrutura de alto desempenho. Trabalhos futuros incluem a execução dos experimentos no Núcleo de Processamento de Alto Desempenho (NPAD) da UFRN e a exploração de outros sinais biomédicos.

Palavras-chave: Redes Neurais Pulsadas; Processamento de sinais biomédicos; Câncer;

TITLE: Use of Spiking Neural Networks in Skin Tumor Image Processing

Abstract

In this work, Spiking Neural Networks (SNNs), a bio-inspired model, were developed and evaluated for skin cancer image classification. The goal of the project is to achieve the high accuracy of Convolutional Neural Networks (CNNs) with significantly lower computational and energy costs. The research employed the pyGeNN and mlGeNN libraries in Google Colab to build and train the SNNs. Model performance was tested on the MNIST, DermaMNIST, and ISIC skin cancer datasets. The network achieved an accuracy of 93.72% on the MNIST test set and 63% on DermaMNIST. For the ISIC images, careful preprocessing was carried out, including resizing, illumination adjustment, and data augmentation, to optimize the images for the network. However, training with the ISIC dataset revealed the limitations of the Google Colab environment, which could not support the computational

demands of the task. The maximum accuracy obtained during training was 65%, highlighting the challenge of training SNNs at scale without specialized hardware. The results suggest that although spiking networks built with GeNN and its libraries show potential, their large-scale application requires significant hyperparameter tuning and, most importantly, the use of high-performance computing infrastructure. Future work includes running the experiments on the High-Performance Computing Center (NPAD) at UFRN and exploring other biomedical signals.

Keywords: Spiking Neural Networks; Biomedical Signal Processing; Cancer

Introdução

Com o avanço da tecnologia e o aumento da disponibilidade de dados biológicos, tem-se observado um crescimento significativo no desenvolvimento de dispositivos vestíveis e da internet das coisas. Para processá-los, as CNNs são eficazes, mas exigem muitos recursos computacionais e energéticos. Em contrapartida, as SNNs são mais eficientes e rápidas, embora ainda pouco utilizadas em sinais biomédicos.

O trabalho de Junayed, Anjum, Sakib e Islam (2021) utiliza imagens de câncer de pele de diferentes repositórios para treinar uma CNN profunda para predição de 4 diferentes classes de câncer (Queratose Actínica, Carcinoma de Células Basais, Melanoma e Carcinoma de Células Escamosas) e obteve uma precisão de 95,98% nos dados de teste. Um resultado que superou a performance de dois modelos pré-treinados, GoogleNet e MobileNet. Esses resultados destacam a eficácia dessa arquitetura para a detecção de tumores de pele.

Em contrapartida, o trabalho de Habaebi et al. (2024) utiliza um modelo de SNNs profunda, Spiking VGG-11, para classificar imagens do ISIC 2019 em melanoma e não melanoma. O modelo foi desenvolvido com o framework SpikingJelly, que é otimizado para simulações de SNNs, com neurônios Leaky Integrate and Fire (LIF) e um PoissonEncoder para converter as entradas em pulsos. O treinamento foi realizado com o otimizador Adaptive Moment Estimation (Adam) e a função de perda Cross Entropy. O modelo, com cerca de 15 milhões de parâmetros, alcançou uma acurácia de 81,3%, resultado competitivo considerando que o modelo é consideravelmente menos complexo que outras arquiteturas, como Spiking VGG-13, que possui aproximadamente 132 milhões de parâmetros.

Este projeto propõe o desenvolvimento de novas arquiteturas de SNNs para obter a mesma precisão das CNNs, mas com menor custo. A abordagem é substituir a arquitetura das CNNs por modelos bioinspirados. Além disso, o projeto usará neurônios com características biofísicas específicas de cada núcleo cerebral.

Os métodos serão aplicados em imagens de tumores de pele para classificá-los como benignos ou malignos. Essa abordagem poderá ser adaptada futuramente para outros sinais biomédicos, como eletromiografia, eletroencefalografia ou ultrassonografia.

O projeto visa criar algoritmos de alta precisão com menores requisitos computacionais e energéticos, facilitando o uso em dispositivos vestíveis e sistemas embarcados e abrindo caminho para o hardware neuromórfico. As novas arquiteturas de SNNs também podem ajudar em estudos sobre o cérebro e o processamento de informações.

Metodologia

Toda implementação da rede foi feita na linguagem Python dentro do ambiente virtual Google Colaboratoy e os códigos estão disponíveis em github.com/cataisout/SNN.

A rede pulsada foi criada fazendo uso do ambiente de simulação GPU-enhanced Neuronal Networks (GeNN) (Yavuz; Turner; Nowotny, 2016), baseado na geração de código para CUDA da Nvidia (NVIDIA et al., 2020). Mais especificamente foram usadas as bibliotecas PyGeNN (Knight; Komissarov; Nowotny, 2021) e mlGeNN (Knight; Nowotny, 2023) que otimizam o processo de criação da rede tornando-a mais eficiente na propagação dos picos de atividade dos neurônios.

As bibliotecas oferecem uma API inspirada no Keras que facilita a criação e o treinamento da rede, sendo particularmente adequada para problemas de aprendizado de máquina que envolvem dimensão temporal (Knight; Nowotny, 2023). A flexibilidade do mlGeNN permite também a fácil prototipagem de diferentes arquiteturas de rede.

Inicialmente, a performance da rede foi testada em conjuntos de dados de referência bem conhecidos, como o MNIST e o DermaMNIST. O principal objetivo dessa fase foi validar a implementação correta das bibliotecas e, ao mesmo tempo, estabelecer uma linha de base para o desempenho em cenários já amplamente estudados para depois testar com imagens mais próximas da realidade médica.

Na segunda etapa, a base de dados utilizada foi a de câncer de pele do ISIC, que contém nove classes distintas de lesões, três de tumores malignos e seis não malignos, sendo a proporção de imagens por categoria de 995 imagens de câncer e 1244 não câncer como mostra a figura 1. Dentro das subcategorias existiam diferentes quantidades disponíveis – o que foi levado em consideração para realizar o balanceamento (Data Augmentation) – e tamanhos, que foram ajustados de forma a preservar a maior quantidade de informações sem comprometer a capacidade de computação da rede. A figura 2 expõe as distribuições dentro de cada classe.

Os dados utilizados passaram por uma série de transformações antes do treinamento. Inicialmente, foi desenvolvida uma pipeline de pré-processamento contemplando transformação para escala de cinza, padronização de brilho, aumento de dados (adição de ruído) além de ajuste de tamanho.

Após essa etapa, as imagens foram carregadas na rede com os parâmetros de treinamento mantidos como os do MNIST para avaliar o desempenho e seguir com mudanças pontuais na arquitetura como número de neurônios por camada, limiares de codificação por latência, número de épocas assim como diferentes tamanhos de imagem.

Ao utilizar dados do ISIC, os experimentos mostraram que treinar SNNs em problemas de grande escala demanda um alto poder computacional e será necessário rodar em ambiente de supercomputação do NPAD. Essa etapa incluiu a análise de dependências, compatibilidades e limitações das ferramentas, permitindo avaliar se eram viáveis para uma infraestrutura de alto desempenho.

Resultados e Discussões

Kight e Nowotny (2023) obtiveram resultados significativos, no conjunto de dados DVS gesture, com a adição de conectividade esparsa com mlGeNN, reduzindo o tempo de

treinamento em aproximadamente 10 vezes, mantendo o desempenho competitivo em relação a outras arquiteturas.

No treinamento com MNIST foi usada uma rede inteiramente feedforward onde cada camada é automaticamente conectada à anterior. Cada pixel é convertido em um único pulso com uma codificação por latência, de acordo com o limiar de escala de cinza definido. Cada pulso é emitido em um ponto no tempo de acordo com o valor de intensidade. Pixels abaixo do limiar são disparados mais cedo e os mais escuros tarde. A figura 3 mostra a fórmula usada no cálculo.

A entrada da rede é em pulsos (SpikeInput) com 256 neurônios, um para cada pixel da imagem (28 x 28), foi usada uma camada escondida com 128 neurônios Leaky Integrate Fire (LIF) com conectividade densa. Na saída foram definidos 10 neurônios Leaky Integrate (LI), com conectividade densa, uma vez que nessa tarefa as imagens possuem 10 classes, a classificação se dá a partir da leitura da soma dos potenciais de membrana por neurônio na saída.

A rede foi compilada utilizando Eprop (Nowotny; Turner; Knight, 2025) e otimizador Adam, um algoritmo usado para treinar do zero que atua como uma aproximação de aprendizado recorrente em tempo real (RTRL). Oferece uma abordagem eficiente baseada em gradiente, uma vez que ele não exige passagem de Backpropagation (BPTT). A rede foi treinada com 10 épocas e a performance foi avaliada com SparseCategoricalAccuracy. No conjunto de treino obteve-se 0.9317 de acurácia e no teste 0.9372.

Os parâmetros foram mantidos para o conjunto DermaMNIST com intuito de avaliar o desempenho da rede para um conjunto mais próximo da realidade. Foi obtida uma acurácia de 0.59 para 10 épocas de treinamento (figura 4) e quando treinado com 40 épocas a acurácia foi de 0.63 (figura 5). Para realizar um treinamento com mais épocas faz-se necessário a execução no Núcleo de Alto Desempenho.

No tratamento das imagens do ISIC, as imagens foram redimensionadas para uma resolução uniforme (600 x 450 pixels, de acordo com a média de tamanho do conjunto). O aumento de dados foi realizado com a aplicação de ruído Gaussiano seguindo uma distribuição normal com média 20 e desvio-padrão 40 e somada aos valores de pixel da imagem. Em seguida os valores foram ajustados para o intervalo [0, 255] para garantir a integridade da imagem, ampliar a variabilidade e mitigar o risco de overfitting. Foram realizados ajustes de iluminação e contraste utilizando a técnica de equalização de histograma adaptativa com limitação de contraste (CLAHE), visando realçar regiões com possíveis padrões importantes para o aprendizado da rede.

O ambiente Colab não suportou manter a proporção de neurônios nas camadas de entrada de acordo com o tamanho das imagens 600 x 450 resultando em falha na RAM. Por isso foi necessário encontrar um novo maior tamanho possível para as imagens (120 x 90). Entre os outros parâmetros alterados de forma exploratória estão, Limiar Theta da codificação por latência, número de neurônios na camada escondida e número de épocas. A figura 6 mostra as diferentes configurações e suas respectivas acurácias. A melhor configuração de acordo com o valor de precisão obtido durante uma época no treinamento foi de 65% no décimo treino. Com 20 épocas, Theta = 0,81 e 758 neurônios na camada escondida. A figura 7 mostra a evolução do treinamento.

Os experimentos, realizados com dados do ISIC, revelaram que o treinamento das SNNs para problemas de grande escala exige um poder computacional considerável, especialmente ao utilizar GPUs. Essa constatação enfatiza a necessidade de buscar soluções que permitam a execução em ambientes com maior capacidade de processamento. Mesmo trabalhando com os parâmetros máximos suportados pelo Google Colab, os resultados foram insatisfatórios, com baixa acurácia, conforme a Figura 6. Isso sugere que, apesar do potencial teórico das SNNs criadas a partir do GeNN, sua aplicação em larga escala ainda demanda um número excessivo de iterações, afinação de hiperparâmetros e hardware neuromórfico.

Conclusão

Embora o pyGeNN e mlGeNN sejam ferramentas promissoras para a prototipagem de Redes Neurais Pulsadas, a sintonização de parâmetros é muito dependente de poder computacional. Diante disso, os próximos passos consistem em utilizar o ambiente de alto desempenho do NPAD para execução de experimentos mais robustos, viabilizando a validação do potencial das bibliotecas. Além disso, faz-se necessário ainda, testar a performance do modelo em sinais de natureza temporal.

Referências

- UNAYED, Masum Shah; ANJUM, Nipa; SAKIB, Abu; ISLAM, Md Baharul. A Deep CNN Model for Skin Cancer Detection and Classification. Csrn, [S.L.], 2021. Západočeská univerzita. <http://dx.doi.org/10.24132/csrn.2021.3101.8>.
- YAVUZ, Esin; TURNER, James; NOWOTNY, Thomas. GeNN: a code generation framework for accelerated brain simulations. Scientific Reports, [S.L.], v. 6, n. 1, 7 Jan. 2016. Springer Science and Business Media LLC. <http://dx.doi.org/10.1038/srep18854>.
- NVIDIA, Vingelmann, P., and Fitzek, F. H. (2020). CUDA, Developer. Available online at: <https://nvidia.com/cuda-toolkit>.
- KNIGHT, James C.; KOMISSAROV, Anton; NOWOTNY, Thomas. PyGeNN: a python library for gpu-enhanced neural networks. Frontiers In Neuroinformatics, [S.L.], v. 15, 22 abr. 2021. Frontiers Media SA. <http://dx.doi.org/10.3389/fninf.2021.659005>.
- KNIGHT, James C; NOWOTNY, Thomas. Easy and efficient spike-based Machine Learning with mlGeNN. Neuro-Inspired Computational Elements Conference, [S.L.], p. 115-120, 11 abr. 2023. ACM. <http://dx.doi.org/10.1145/3584954.3585001>.
- Tutorials — mlGeNN documentation. Disponível em: <<https://ml-genn.readthedocs.io/en/latest/tutorials/index.html>>. Acesso em: 31 ago. 2025.
- NOWOTNY, Thomas; TURNER, James P; KNIGHT, James C. Loss shaping enhances exact gradient learning with Eventprop in spiking neural networks. Neuromorphic Computing And Engineering, [S.L.], v. 5, n. 1, p. 014001, 21 Jan. 2025. IOP Publishing. <http://dx.doi.org/10.1088/2634-4386/ada852>.
- HABAEBI, Mohamed Hadi; ZULKARNAIN, Muhammad Amin; GUNAWAN, Teddy Surya; FITRIAWAN, Helmy; KARTIWI, Mira; ABD.RAHMAN, Faridah. Development of

Skin Cancer Detection Using Deep Spiking Neural Networks. 2024 Ieee 10Th International Conference On Smart Instrumentation, Measurement And Applications (Icsima), [S.L.], p. 125-129, 30 jul. 2024. IEEE. <http://dx.doi.org/10.1109/icsima62563.2024.10675560>.

Anexos

Figura 1

Figura 2

Figura 3

Figura 4

Figura 5

Figura 6

Figura 7

CÓDIGO: ET1630

AUTOR: HOSANA IASMIN CASTRO DOS SANTOS

COAUTOR: JADY LIMA DA SILVA

COAUTOR: GABRIEL ROCHA DOS SANTOS

COAUTOR: RICARDO JOSE MATOS DE CARVALHO

ORIENTADOR: PATRICK CESAR ALVES TERREMATTE

TÍTULO: Aplicação de Modelagem de Tópicos com BERTopic para Análise de Relatos de Desastres do Sistema Web NOAH

Resumo

O Gerenciamento de Riscos e Desastres (GRD) demanda a análise rápida de grandes volumes de relatos textuais da população. Este estudo aplicado e exploratório detalha o desenvolvimento do módulo de Processamento de Linguagem Natural (PLN) Noah-IA, integrado ao sistema Noah, para automatizar a triagem de ocorrências em Natal (RN). A metodologia partiu da criação de um corpus sintético com ChatGPT, e empregou o modelo BERTopic para a descoberta não supervisionada de tópicos e a biblioteca spaCy para a classificação da temporalidade dos relatos (risco vs. pós-fato). Os resultados demonstram a capacidade do modelo em identificar não apenas temas, mas "cenários de desastre" (e.g., queda de árvore sobre postes), validando a relevância dos achados ao compará-los com dados oficiais da Defesa Civil. Conclui-se que a abordagem proposta é eficaz para transformar dados não-estruturados em informação acionável, com potencial para otimizar a priorização de respostas a emergências.

Palavras-chave: Processamento de Linguagem Natural. Gerenciamento de Desastres.

TITLE: Topic Modeling Application with BERTopic for Analysis of Disaster Reports from the NOAH Web System

Abstract

Disaster Risk Management (DRM) requires the rapid analysis of large volumes of textual reports from the population. This applied and exploratory study details the development of the Natural Language Processing (NLP) module noah-ia, integrated into the Noah system, to automate the triage of incidents in Natal, Brazil. The methodology began with the creation of a synthetic corpus using ChatGPT, and employed the BERTopic model for unsupervised topic discovery and the spaCy library for classifying the reports' temporality (risk vs. after-the-fact). The results demonstrate the model's ability to identify not just themes, but complex

"disaster scenarios" (e.g., trees falling on power lines), validating the findings by comparing them with official Civil Defense data. The study concludes that the proposed approach is effective in transforming unstructured data into actionable information, with the potential to optimize the prioritization of emergency responses, contributing to urban resilience.

Keywords: Natural Language Processing. Disaster Management.

Introdução

O Gerenciamento de Riscos e Desastres (GRD) tem se tornado uma área de crescente importância, impulsionada pela intensificação de eventos climáticos e pela crescente vulnerabilidade dos centros urbanos. Neste cenário, a comunicação eficaz entre os órgãos públicos, como a Defesa Civil, e a população é um pilar fundamental para uma resposta rápida e coordenada. A relevância deste tema para Natal-RN é evidenciada pela recente apresentação do Plano Municipal de Redução de Risco (PMRR), um esforço que mapeou áreas vulneráveis a deslizamentos e alagamentos e reforçou a necessidade de uma gestão de riscos participativa (**PREFEITURA DO NATAL, 2025**). Contudo, a troca de informações durante a fase de resposta a desastres permanece como um dos maiores desafios operacionais, muitas vezes dificultada por ruídos de comunicação e pela falta de canais eficientes para a coleta de dados em tempo real, um desafio amplamente documentado na literatura da área (**IMRAN et al., 2015**).

Visando endereçar este desafio no contexto da cidade de Natal-RN, o sistema Noah surge como uma plataforma colaborativa que utiliza a Inteligência Artificial (IA) para apoiar o GRD. Através de um sistema web e de um chatbot, o Noah propõe um canal direto para que a população possa notificar riscos e incidentes, ajudando a proteger a comunidade local. No entanto, ao abrir essa via de comunicação, o sistema se depara com um desafio tecnológico fundamental: o recebimento de um grande volume de dados em formato de texto livre. Estes relatos, expressos em linguagem natural, são heterogêneos, informais e não estruturados, tornando a triagem e a priorização manual uma tarefa inviável e lenta, que pode comprometer o tempo de resposta a eventos críticos.

Para analisar os relatos da população e gerar informações que ajudem na tomada de decisão, este trabalho apresenta a criação de um módulo de Processamento de Linguagem Natural (PLN) para o sistema Noah. O objetivo principal deste artigo é, portanto, apresentar o desenvolvimento e a aplicação do modelo BERTopic para a categorização automática de relatos de desastres no chatbot do Noah.

Metodologia

Um desafio primário para o desenvolvimento de sistemas de IA neste domínio é a disponibilidade de dados de treinamento. Para superar a falta de um corpus público específico para desastres no contexto brasileiro, o ponto de partida metodológico foi a criação de um conjunto de dados sintético. Especificamente, foi utilizado o modelo de linguagem generativa ChatGPT para criar um volume de 206 exemplos textuais que simulam relatos de cidadãos para diversas categorias de ocorrências como

alagamento, enxurrada, deslizamento, incêndio, etc. Os textos gerados foram então copiados em um arquivo de formato CSV(Comma-Separated Values), que serviu como fonte de entrada para a etapa seguinte.

A preparação deste corpus para a modelagem envolveu um processo de pré-processamento textual, resultando na normalização dos dados através da conversão para minúsculas, garantindo que palavras como “Chuva” e “chuva” fossem tratadas da mesma forma, remoção de pontuações, consistindo na retirada de todos os caracteres que não eram letras para focar apenas no conteúdo do texto e eliminação de stopwords, que são preposições muito comuns em português que não carregam um significado importante para definir um tópico (como “de”, “a”, “o”, “que”, “em”).

Para a modelagem de tópicos, foi utilizado o modelo BERTopic (**GROOTENDORST, 2022**), que é uma abordagem moderna baseada em embeddings contextuais, diferente da abordagem clássica e mais difundida, a Alocação Latente de Dirichlet (LDA) (**BLEI; NG; JORDAN, 2003**).

Essa distinção é fundamental: o LDA, trata os documentos como uma mistura de tópicos e os tópicos como uma distribuição de palavras, o que o torna dependente da frequência delas, mas cego ao contexto semântico (**EGGER; YU, 2022**). Em contraste, o BERTopic fundamenta-se na hipótese de que documentos com significados similares devem ser agrupados. O resultado é um modelo mais robusto, que não exige a definição prévia do número de tópicos e demonstra performance superior na análise de textos, sendo, portanto, mais adequado às necessidades deste módulo do sistema NOAH.

O processo inicia-se com a etapa de vetorização, na qual cada relato de desastre é convertido em um embedding semântico. Para esta tarefa, foi utilizado o modelo SentenceTransformer com a arquitetura pré -treinada ‘paraphrase-multilingual-mpnet-base-v2’, escolhida pela sua eficácia em capturar nuances contextuais em português.

Dado que os embeddings semânticos gerados são de alta dimensionalidade, foi aplicado o algoritmo UMAP (**MCINNES; HEALY; MELVILLE, 2018**) para a redução e projeção dos dados. O principal parâmetro, n_neighbors, foi configurado com o valor 5. Esta escolha se justifica pela escala do conjunto de dados (206 amostras), priorizando a formação de clusters mais bem definidos. O parâmetro n_components foi definido como 2, instruindo o algoritmo a projetar os dados em um espaço bidimensional, o que permite a posterior visualização gráfica dos resultados.

Para o cálculo de distância entre os pontos no espaço, foi selecionada a métrica 'cosine'. Diferente de métricas espaciais, a similaridade de cosseno não mede a magnitude dos vetores, mas sim o ângulo entre eles. Essa abordagem é particularmente eficaz para dados textuais, pois permite identificar a similaridade de significado entre dois relatos, independentemente da quantidade de palavras em cada um. Por fim, o parâmetro random_state foi fixado em 42, uma medida essencial para garantir a reprodutibilidade dos resultados em futuras execuções do código.

Subsequentemente, foi aplicado o algoritmo de clusterização HDBSCAN nos vetores de dimensionalidade reduzida. Seus parâmetros foram ajustados para `min_samples` com valor 3 e `min_cluster_size` com valor 4. Esta configuração específica instrui o modelo a ser mais rigoroso na formação de clusters, permitindo que um tópico seja formalmente identificado a partir de apenas quatro relatos distintos. Por fim, para garantir a interpretabilidade dos clusters formados, foi utilizado um vocabulário pré-definido. Este vocabulário foi elaborado para incluir termos específicos e relevantes ao domínio de desastres, abrangendo categorias como 'Hidrológicos' (e.g., inundação, alagamento), 'Geológicos' (e.g., deslizamento), 'Meteorológicos' (e.g., tempestade) e 'Riscos Estruturais' (e.g., poste, árvore). A utilização deste vocabulário direciona o modelo a focar nas palavras mais pertinentes ao problema, assegurando que a representação final de cada tópico fosse concisa.

Para enriquecer a análise uma segunda classificação foi adicionada ao pipeline. Esta camada, desenvolvida com a biblioteca spaCy, realiza uma análise morfológica para determinar a temporalidade dos relatos. A função implementada examina os verbos em cada texto para diferenciar entre eventos consumados ("pós-fato") e situações de perigo em andamento ("risco"). O resultado final do pipeline é, portanto, um dado categorizado tematicamente pelo BERTopic e classificado quanto à sua urgência pela análise temporal.

Por fim, o módulo de PLN desenvolvido não opera de forma isolada, mas sim como um componente de serviço dentro da arquitetura integrada do sistema Noah, disponível em <https://noah.imd.ufrn.br/>, atuando como o núcleo que interpreta as mensagens dos cidadãos. O fluxo operacional para o registro de uma nova ocorrência foi projetado da seguinte forma:

Cidadão → WhatsApp → NOAH-Bot (que consome o) → [Módulo de PLN] → API → Plataforma Noah Web → Banco de Dados

Dessa forma, o módulo de PLN atua como um microsserviço permitindo que todo o ecossistema Noah opere com informações já classificadas.

Resultados e Discussões

Identificação e Análise dos Tópicos de Ocorrências

O resultado inicial da aplicação do modelo BERTopic foi a identificação de um conjunto de tópicos temáticos a partir dos relatos. A Figura 1 (Anexos) expõe os termos de maior relevância para uma amostra dos tópicos descobertos.

Por exemplo, o Tópico 2 representa o cenário de "árvore", "tempestade" e "poste".

A discussão em torno deste resultado inicial reside na capacidade do modelo de, autonomamente, transformar relatos livres e não-estruturados dos cidadãos em dados estruturados e categorizados. Este é o primeiro passo para otimizar a integração dessas informações com outras fontes e para automatizar a triagem de

ocorrências. A variedade de tópicos identificados, demonstra a abrangência da análise.

Estrutura Semântica e a Descoberta de Cenários

Foi analisada também a relação semântica entre os tópicos. O dendrograma de clusterização hierárquica (Figura 2, Anexos) ilustra a estrutura de aglomeração, revelando a formação de super-grupos lógicos, como o que une ocorrências de "incêndio" e "explosão" e outro que agrupa "deslizamento" e "enxurrada". De forma complementar, a matriz de similaridade (Figura 3, Anexos) quantifica essa relação, com tons mais escuros indicando maior proximidade de significado entre os tópicos.

O ponto mais relevante desta análise estrutural é a constatação de que o modelo não identificou apenas palavras-chave, mas sim "cenários de desastre". O Tópico 2 (árvore_tempestade_poste), por exemplo, não representa três temas distintos, mas provavelmente um evento único de "queda de árvore sobre a rede elétrica durante uma tempestade". Esta capacidade de extrair cenários é uma vantagem de modelos contextuais como o BERTopic sobre abordagens clássicas. Para a Defesa Civil, um relato categorizado com este cenário é muito mais informativo e direciona uma resposta mais eficaz do que um relato genericamente classificado como "árvore" ou "tempestade".

Validação Funcional e Temporalidade

Para testar se o sistema funcionava na prática, foi realizado um teste local com o seguinte relato: "estou dirigindo e acabei de ver que abriu uma cratera enorme próximo ao natal shopping, alguns motoqueiros quase caíram". O sistema retornou duas classificações para essa frase: o tópico foi identificado como "cratera poste" e a temporalidade como "pós-fato".

A análise desses resultados mostrou dois lados. A classificação da temporalidade como "pós-fato" estava correta e se mostrou muito útil, pois saber a diferença entre um "risco" (algo acontecendo agora) e um "pós-fato" é fundamental para a Defesa Civil priorizar suas ações.

Por outro lado, a classificação do tópico como "cratera poste" foi um acerto parcial. O modelo acertou o tema principal ("cratera"), mas a inclusão de "poste" foi uma imprecisão, já que a palavra não estava na mensagem. Isso provavelmente aconteceu pela forma como as palavras foram associadas nos dados de treinamento. O resultado evidencia que, para trabalhos futuros, é necessário treinar o modelo com um volume de dados maior e mais variado para refinar a precisão dos tópicos.

Comparação com Dados Oficiais e Relevância Prática

A validação final da relevância do modelo foi realizada comparando os tópicos gerados com os dados oficiais de ocorrências de 2024, disponibilizados na plataforma Noah Web (Figura 4, Anexos). A análise comparativa revelou uma certa correspondência. O tópico de maior frequência identificado pelo modelo, "Inundação e Alagamento", alinha-se diretamente com a categoria oficial mais recorrente de 2024, "Transbordamento de Lagoa".

Esta correspondência demonstra que o modelo de PLN foi capaz de, a partir de texto livre e de forma não supervisionada, reconstruir eventos similares àqueles utilizados pelos próprios agentes da Defesa Civil. Isso não apenas valida a eficácia da abordagem metodológica, mas confirma o potencial da ferramenta para automatizar a triagem de relatos com alinhamento à realidade operacional do órgão.

Conclusão

Os resultados mostraram que a abordagem funcionou bem. A maior descoberta foi que o modelo aprendeu a identificar "cenários" completos (como 'queda de árvore em poste'), o que é muito mais útil do que apenas encontrar palavras-chave soltas. Além disso, a comparação entre os tópicos encontrados pelo sistema e as categorias oficiais da Defesa Civil de Natal revelou uma alta correspondência, o que prova a relevância da ferramenta e sua capacidade de entender os problemas reais da cidade.

A principal contribuição deste projeto é um protótipo funcional que mostra ser possível automatizar a triagem de relatos de desastres. Na prática, isso pode ajudar a Defesa Civil a otimizar seu tempo e recursos, respondendo mais rápido às emergências mais urgentes. O trabalho serve como um exemplo prático de como a IA pode ser aplicada para ajudar a tornar as cidades mais seguras.

Apesar dos bons resultados, foi identificada uma limitação importante: o tempo de resposta do modelo, que nos testes realizados foi alto para uma aplicação em tempo real. Por isso, como sugestões para trabalhos futuros, recomenda-se: otimizar o modelo para que se torne mais rápido, usando técnicas específicas e equipamentos mais potentes como placas de vídeo (GPUs); treinar o modelo com os dados reais que o NOAH-Bot irá coletar, para que fique ainda mais inteligente e adaptado à forma como as pessoas de Natal falam e reconfigurar o modelo para melhor precisão de tópicos.

Referências

PREFEITURA DO NATAL. Prefeitura apresenta Plano de Redução de Risco e abre consulta pública para contribuições. Natal, 2025. Disponível em: <https://www.natal.rn.gov.br/news/post2/43137>. Acesso em: 29 ago. 2025.

IMRAN, M.; CASTILLO, C.; DIAZ, F.; MEIER, P. Processing social media messages in crises: A survey. *ACM Computing Surveys*, v. 47, n. 4, p. 1-38, 2015.

GROOTENDORST, M. BERTopic: Neural topic modeling with a class-based TF-IDF procedure. *arXiv preprint arXiv:2203.05794*, 2022

BLEI, D. M.; NG, A. Y.; JORDAN, M. I. Latent Dirichlet Allocation. *Journal of Machine Learning Research*, v. 3, p. 993-1022, 2003.

EGGER, R.; YU, J. A topic modeling comparison between LDA, NMF, Top2Vec, and BERTopic. In: *Frontiers in Artificial Intelligence and Applications - Volume 349: Knowledge Management in Organizations*. IOS Press, 2022. p. 135-147.

MCINNES, L.; HEALY, J.; MELVILLE, J. UMAP: Uniform Manifold Approximation and Projection for Dimension Reduction. arXiv preprint arXiv:1802.03426, 2018.

Anexos

Figura 1 - Gráfico de barras com as pontuações c-TF-IDF das palavras-chave mais representativas para uma amostra dos tópicos identificados.

Figura 2 - Dendrograma da clusterização hierárquica ilustrando a relação de similaridade semântica e a formação de super-grupos entre os tópicos.

Figura 3 - Matriz de similaridade de cosseno entre os tópicos, onde tons mais escuros indicam maior proximidade semântica.

Figura 4 - Distribuição dos tipos de ocorrências registrados pela Defesa Civil de Natal no ano de 2024.

CÓDIGO: ET1637

AUTOR: RAMON CÂNDIDO JALES DE BARROS

ORIENTADOR: PATRICK CESAR ALVES TERREMATTE

TÍTULO: Especificação de teorias formais para matemática discreta via programação funcional e assistentes de demonstração

Resumo

Este projeto de Iniciação Científica visa desenvolver um método para a geração automática de conjecturas em matemática discreta com níveis de complexidade controlados. Utilizando programação funcional e o assistente de demonstração Coq, o projeto implementará teorias formais para temas como teoria dos conjuntos, relações de ordem, funções e teoria dos números. A metodologia se baseia na criação de "esquemas de prova", que são demonstrações mínimas, para então buscar ou gerar novas conjecturas que possuam uma estrutura de prova isomórfica, garantindo assim uma complexidade similar. A pesquisa investigará a teoria da complexidade de demonstrações, explorando a construção de árvores de derivação mínimas. Como aplicação prática, o projeto contribuirá para o desenvolvimento de uma ferramenta computacional educacional capaz de gerar exercícios de demonstração personalizados, otimizando o trabalho de educadores e auxiliando no processo de aprendizagem dos alunos.

Palavras-chave: Matemática Discreta, Assistentes de Prova, Programação Funcional, Coq.

TITLE: The Implementation of Formal Theories for Discrete Mathematics via Functional Programming and Proof Assistants.

Abstract

This Scientific Initiation project aims to develop a method for the automatic generation of conjectures in discrete mathematics with controlled complexity levels. Utilizing functional programming and the Coq proof assistant, the project will implement formal theories for topics such as set theory, order relations, functions, and number theory. The methodology is based on the creation of "proof schemes," which are minimal proofs, to then search for or generate new conjectures that possess an isomorphic proof structure, thereby ensuring similar complexity. The research will investigate proof complexity theory, exploring the construction of minimal derivation trees. As a practical application, the project will contribute to the development of an educational software tool capable of generating personalized proof exercises, optimizing the work of educators and assisting in the students' learning process.

Keywords: Discrete Mathematics, Proof Assistants, Functional Programming, Coq.

Introdução

A criação de exercícios e avaliações em disciplinas de exatas, como a Matemática Discreta, representa um desafio significativo para educadores, que despendem um tempo considerável

na elaboração manual de questões. Embora a Geração Automática de Questões (AQG) tenha surgido como uma solução promissora para otimizar esse processo, a maioria das técnicas existentes falha em oferecer um controle preciso sobre o nível de dificuldade dos problemas gerados. Essa lacuna pode resultar em avaliações desiguais, onde alguns alunos enfrentam questões com complexidade lógica muito superior à de seus colegas, prejudicando a isonomia do processo de aprendizagem e avaliação. Nosso projeto propõe uma abordagem inovadora para resolver essa questão, fundamentada na teoria da prova. A ideia central é desenvolver um método que gere conjecturas matemáticas com complexidade controlada a partir do que chamamos de "esquemas de prova". Um esquema de prova é, essencialmente, a estrutura de uma demonstração mínima. Partindo do princípio de que exercícios que compartilham uma estrutura de prova isomórfica (ou seja, seguem os mesmos passos lógicos) possuem um nível de dificuldade similar, usaremos programação funcional e o assistente de provas Coq para implementar as teorias formais da matemática discreta e buscar automaticamente por conjecturas que se encaixem nesses esquemas. Ao implementar esse método, o projeto terá uma contribuição tanto teórica quanto prática. No campo teórico, avançaremos nos estudos de complexidade de demonstrações e especificação formal. Na prática, o objetivo final é desenvolver uma ferramenta computacional com interface web capaz de automatizar a geração de listas de exercícios personalizados. Tal ferramenta poderá auxiliar educadores a economizar tempo e a criar avaliações mais justas e eficazes, além de fornecer aos alunos um recurso de estudo adaptativo e alinhado ao seu progresso na disciplina.

Metodologia

A abordagem metodológica deste projeto centrou-se na utilização do assistente de provas Coq e de técnicas de programação funcional para implementar as teorias formais da matemática discreta, abrangendo tópicos como operações sobre conjuntos, relações de ordem, funções e teoria dos números. O núcleo do trabalho consistiu em definir a complexidade de uma demonstração pela quantidade de regras da teoria aplicadas, permitindo assim a construção de provas mínimas, que chamamos de "esquemas de prova". Uma vez estabelecido um esquema, desenvolvemos um método para buscar e gerar novas conjecturas cujas árvores de derivação fossem isomórficas ao esquema original, garantindo que os novos exercícios tivessem um nível de complexidade similar ao da conjectura de partida. Para colocar o método em prática, o projeto envolveu a implementação de uma ferramenta computacional que automatiza todo o processo de geração. Essa ferramenta foi desenvolvida para receber um esquema de prova como entrada e, a partir dele, realizar a busca por conjecturas isomórficas, resultando na criação de exercícios de demonstração personalizados. Durante o desenvolvimento, também investigamos a literatura sobre cut-based tableaux como forma de mecanizar e otimizar a construção das árvores de derivação mínimas. O objetivo final do processo foi, portanto, criar um sistema funcional para a geração automática de questões, visando reduzir o esforço manual de educadores na criação de avaliações.

Resultados e Discussões

Resultados e Discussões Como resultado prático deste projeto, desenvolvemos e aplicamos um protótipo de ferramenta para a geração automática de exercícios em uma turma de Fundamentos Matemáticos da Computação II (FMC2). No experimento, cada aluno recebeu três asserções individualizadas sobre teoria dos conjuntos para provar ou refutar no assistente de provas Coq. O principal resultado observado foi uma intensa demanda por suporte da

equipe de monitoria, com um alto volume de interações remotas, o que evidencia que a ferramenta cumpriu o objetivo de gerar atividades que engajaram os alunos. Ao mesmo tempo, ficou claro que a resolução de provas formais em um ambiente baseado em texto como o Coq possui uma curva de aprendizado acentuada, tornando o auxílio humano um componente fundamental para o sucesso da atividade nesta fase. A análise desses resultados preliminares aponta para uma validação parcial dos objetivos do projeto. A capacidade de gerar e distribuir exercícios únicos em larga escala vai ao encontro do que a literatura propõe sobre avaliação personalizada para promover o aprendizado e mitigar o plágio (MANOHARAN, 2017). A experiência contrastou com o uso de ferramentas mais gráficas, como o Edukera, empregadas em semestres anteriores, sugerindo que a maior expressividade lógica e o rigor do Coq vêm acompanhados de uma maior exigência de abstração por parte do aluno. O engajamento observado, embora positivo, reforça a percepção de que a ferramenta, em seu estado atual, funciona melhor como um complemento ao ensino, mediado pela figura do monitor, do que como uma plataforma de aprendizado totalmente autônoma. É crucial ressaltar que, por o trabalho não estar finalizado, estes são resultados parciais. A principal hipótese do projeto — de que nosso método gera exercícios com níveis de complexidade controlados e similares — ainda precisa ser formalmente validada através de uma análise aprofundada das resoluções submetidas pelos alunos. As limitações do experimento incluem sua aplicação em uma única turma e a cobertura de um tópico restrito da disciplina. Os próximos passos, portanto, são finalizar a análise dos dados para validar nossa métrica de complexidade, expandir o gerador para cobrir outros conteúdos de FMC e aprimorar a ferramenta com mecanismos de feedback que possam, futuramente, reduzir a dependência do suporte humano.

Conclusão

Este projeto de Iniciação Científica se propôs a enfrentar o desafio da criação de avaliações personalizadas em matemática discreta e, ao longo deste ano, logrou desenvolver um método inovador baseado em esquemas de prova, implementando um protótipo funcional para a geração automática de exercícios no assistente de provas Coq. A aplicação deste protótipo em um ambiente real de sala de aula, no contexto da disciplina de Fundamentos Matemáticos da Computação, validou a viabilidade prática da abordagem e demonstrou um potencial significativo para aumentar o engajamento dos alunos com a matéria, como evidenciado pela alta demanda por suporte da monitoria. Embora os resultados preliminares sejam promissores, a validação formal da hipótese de que os exercícios gerados possuem um nível de complexidade efetivamente controlado permanece como a principal tarefa a ser concluída. Ainda assim, o trabalho estabelece uma base sólida, contribuindo com uma metodologia e uma ferramenta que unem a teoria da computação e a educação, e abre caminhos claros para futuras pesquisas, incluindo a expansão do sistema para outros domínios da matemática e o aprimoramento dos mecanismos de feedback para a construção de uma plataforma de aprendizado mais autônoma.

Referências

D'Agostino, M. Tableau Methods for Classical Propositional Logic. In: Formal Specifications, Automated Reasoning, The use of Computers in Education. Dordrecht: Springer Netherlands, 1999. p. 45–123. KURDI, G.; LEO, J.; PARSIA, B.; SATTler, U.; AL-EMARI, S. A systematic review of automatic question generation for educational

purposes. *International Journal of Artificial Intelligence in Education*, v. 30, n. 1, p. 121–204, 2020. MANOHARAN, S. Personalized assessment as a means to mitigate plagiarism. *IEEE Transactions on Education*, v. 60, n. 2, p. 112–119, 2017. MILLAR, R.; MANOHARAN, S. Repeat individualized assessment using isomorphic questions: a novel approach to increase peer discussion and learning. *International Journal of Educational Technology in Higher Education*, v. 18, n. 1, p. 22, 2021. MENDES, João; MARCOS, João. A method for automated generation of exercises with similar level of complexity. In: *WORKSHOP BRASILEIRO DE LÓGICA (WBL)*, 4. , 2023, João Pessoa/PB. Anais [...]. Porto Alegre: Sociedade Brasileira de Computação, 2023 . p. 17-24. ISSN 2763-8731. DOI: <https://doi.org/10.5753/wbl.2023.230660>. PIERCE, Benjamin C. et al. *Logical Foundations*. Vol. 1. Electronic textbook. 2023. Disponível em: <https://softwarefoundations.cis.upenn.edu/lf-current/index.html>. Acesso em: 7 jul. 2024. SINGHAL, R.; GOYAL, S.; HENZ, M. User-defined difficulty levels for automated question generation. In: *2016 IEEE 28th International Conference on Tools with Artificial Intelligence (ICTAI)*. p. 828–835, 2016. WANG, K.; SU, Z. Dimensionally guided synthesis of mathematical word problems. In: *Proceedings of the Twenty-Fifth International Joint Conference on Artificial Intelligence, IJCAI'16*. AAAI Press, 2016. p. 2661–2668.

CÓDIGO: ET1649

AUTOR: JOSE EDUARDO MONTEIRO DOS SANTOS

COAUTOR: RAMON CÂNDIDO JALES DE BARROS

COAUTOR: JOÃO LUCAS DE MORAES PEREIRA

COAUTOR: DAVI NASCIMENTO GOMES DE PAIVA

COAUTOR: JEREMIAS PINHEIRO DE ARAUJO ANDRADE

ORIENTADOR: PATRICK CESAR ALVES TERREMATTE

TÍTULO: Especificação de teorias formais para Aritmética via programação funcional e assistentes de demonstração

Resumo

A resolução de exercícios é fundamental para o aprendizado e a avaliação em disciplinas de lógica e matemática. No entanto, a elaboração de listas de exercícios personalizadas para turmas numerosas é um desafio para os docentes, dificultando o acompanhamento individualizado e abrindo espaço para desonestidade acadêmica. Este trabalho apresenta o GREAT (Generator of Propositional Logic Exercises), um gerador de exercícios que produz conjecturas em lógica proposicional para tarefas de prova em Dedução Natural e refutação por contraexemplo. A ferramenta foi desenvolvida para apoiar o ensino de fundamentos matemáticos para a computação, permitindo que estudantes pratiquem a construção de provas formais com o auxílio do assistente de prova interativo Rocq. O sistema oferece ampla customização, permitindo controlar o número de átomos, conectivos lógicos, complexidade das fórmulas e a quantidade de premissas. Recentemente, o GREAT passou por melhorias significativas, incluindo a otimização de performance com um novo SAT solver, uma interface web redesenhada para melhor usabilidade e a capacidade de exportar exercícios em formato LaTeX. Disponibilizado como software livre, o GREAT se consolida como uma ferramenta pedagógica robusta para enriquecer o ensino de lógica e engajar estudantes em tarefas formais.

Palavras-chave: Programação Funcional, Matemática discreta, Assistente de demonstração

TITLE: Specification of Formal Theories for Arithmetic via Functional Programming and Proof Assistants

Abstract

Solving exercises is crucial for learning and assessment in logic and mathematics. However, creating personalized exercise lists for large classes is a challenge for instructors, hindering individualized monitoring and creating opportunities for academic dishonesty. This paper presents GREAT (Generator of Propositional Logic Exercises), an exercise generator that produces propositional logic conjectures for Natural Deduction proof tasks and refutation by counterexample. The tool was developed to support the teaching of mathematical foundations

for computer science, allowing students to practice formal proof construction with the interactive proof assistant Rocq. The system offers extensive customization, allowing control over the number of atoms, logical connectives, formula complexity, and number of premises. Recently, GREAT has undergone significant improvements, including performance optimization with a new SAT solver, a redesigned web interface for better usability, and the ability to export exercises in LaTeX format. Available as free software, GREAT stands as a robust pedagogical tool to enrich logic education and engage students in formal tasks.

Keywords: Functional Programming, Discrete Mathematics, Proof Assistant

Introdução

A fixação de conteúdo e a avaliação do aprendizado em disciplinas de ciências exatas, como a Matemática Discreta, dependem fortemente da resolução de exercícios. Em turmas numerosas e com níveis de conhecimento heterogêneos, a criação de atividades individualizadas e personalizadas torna-se uma tarefa complexa para os docentes. Esta dificuldade não apenas impede um acompanhamento pedagógico mais preciso, mas também favorece a ocorrência de plágio e o uso indiscriminado de ferramentas de inteligência artificial generativa sem a devida reflexão por parte dos alunos.

Nesse contexto, ferramentas de Geração Automática de Questões (AQG, do inglês Automatic Question Generation) surgem como uma solução viável para auxiliar professores na produção massiva de exercícios únicos, garantindo que cada estudante receba uma tarefa distinta e adequada ao seu nível de conhecimento.

Este trabalho apresenta o GREAT (Gerador Automático de Exercícios de Lógica Formal), um gerador de exercícios desenvolvido na UFRN como ferramenta de apoio ao ensino de lógica em cursos de computação. O sistema é projetado para criar conjecturas em lógica proposicional, que podem ser utilizadas em tarefas de demonstração, utilizando o sistema de Dedução Natural, ou de refutação por meio da apresentação de um contraexemplo. Atualmente, o GREAT é empregado em conjunto com o assistente de prova interativo Rocq, onde os alunos praticam a escrita de provas formais a partir dos exercícios gerados e a plataforma ProofWeb, que permite que os exercícios e as turmas sejam acompanhadas de forma on-line.

O objetivo deste relatório é apresentar as atividades desenvolvidas no último ano, que resultaram em melhorias significativas na funcionalidade, usabilidade e desempenho do GREAT. As atualizações transformaram a ferramenta em uma aplicação web de código aberto, com o intuito de distribuí-la livremente para que outros professores e instituições possam utilizá-la como um recurso pedagógico para fortalecer o desenvolvimento do raciocínio lógico e formal dos estudantes.

Metodologia

O GREAT foi desenvolvido como uma aplicação web utilizando PHP e JavaScript, com foco em modularidade e personalização. O método para a geração de exercícios consiste em um processo parametrizado que acaba na verificação de validade da conjectura por meio de um solucionador SAT (Satisfiability Problem).

O processo inicia-se com a configuração dos parâmetros pelo usuário através da interface web. O docente pode definir um parâmetro sobre a complexidade e a estrutura dos exercícios. Os principais parâmetros de customização incluem:

Átomos Proposicionais: Definição do número e dos nomes das variáveis proposicionais (ex: A, B, C).

Conectivos Lógicos: Seleção dos operadores a serem incluídos (negação, conjunção, disjunção, implicação, bi-implicação).

Complexidade da Proposição: Controle da dificuldade por meio da definição do número de conectivos em cada fórmula.

Número de Premissas: Definição da quantidade de premissas que compõem o argumento.

Para garantir a relevância pedagógica dos exercícios, foram adicionados requisitos opcionais:

Relevância: Cada premissa deve compartilhar ao menos um átomo proposicional com a conclusão, evitando premissas desconexas.

Necessidade: Todas as premissas devem ser estritamente necessárias para a prova, eliminando informações supérfluas.

Contingência: A conjunção das premissas não pode ser uma tautologia nem uma contradição.

Tipo de Conjectura: O sistema pode ser configurado para gerar apenas conjecturas prováveis, apenas refutáveis ou uma mistura aleatória.

Com os parâmetros definidos, o GREAT gera aleatoriamente um conjunto de proposições para servirem como premissas e conclusão. Em seguida, o sistema verifica se o conjunto atende aos requisitos de complexidade e relevância. Por fim, a validade da conjectura (i.e., se as premissas implicam a conclusão) é determinada por um SAT solver. Se a implicação ($\text{premissa1} \wedge \text{premissa2} \wedge \dots \rightarrow \text{conclusão}$) for uma tautologia, a conjectura é classificada como "demonstrável". Caso contrário, é classificada como "refutável".

Nos últimos meses, o módulo de verificação foi otimizado com a substituição do SAT solver anterior pelo Picosat, um solucionador mais eficiente. Adicionalmente, implementou-se uma técnica de memoização, que armazena resultados de cálculos já realizados para acelerar a geração de exercícios complexos ou em grande volume, evitando recomputações desnecessárias. A ferramenta também foi aprimorada para gerar saídas em múltiplos formatos, incluindo Unicode, LaTeX e código para o assistente de provas Rocq.

Resultados e Discussões

As atividades de pesquisa e desenvolvimento realizadas no último ano resultaram em melhorias substanciais para o GREAT, consolidando-o como uma ferramenta pedagógica mais robusta, eficiente e acessível. Os principais resultados podem ser divididos em três áreas: funcionalidade, usabilidade e desempenho.

1. Suporte à Exportação em LaTeX e Múltiplos Formatos: Uma das melhorias mais significativas foi a implementação de um sistema de exportação de exercícios para o formato LaTeX. Anteriormente, a saída era limitada a texto simples (Unicode) e ao formato do Rocq.

Agora, os professores podem gerar listas de exercícios formatadas profissionalmente de maneira automática, utilizando templates LaTeX customizáveis. Este recurso economiza horas de trabalho manual na preparação de materiais didáticos e, crucialmente, permite a geração de variantes únicas de um mesmo exercício, combatendo o plágio e a troca de respostas entre alunos. A Figura 3 exemplifica uma saída em LaTeX, demonstrando a clareza e o padrão editorial alcançados.

2. Interface de Usuário Redesenhada: A interface web do GREAT foi completamente redesenhada com foco na experiência do usuário (UX). O novo design (Figura 1) é mais limpo, intuitivo e simplificado. As opções de customização, antes dispostas de forma densa, agora estão organizadas logicamente, guiando o usuário através do processo de configuração dos exercícios. Essa melhoria diminui a curva de aprendizado da ferramenta e torna seu uso mais agradável, incentivando a adoção por parte de educadores que não possuem grande familiaridade com ferramentas de linha de comando.

3. Otimização de Desempenho: A eficiência do processo de geração de exercícios era um gargalo, especialmente para problemas com alta complexidade (muitos conectivos e átomos). Para resolver isso, o SAT solver legado foi substituído pelo Picosat, uma implementação moderna e altamente otimizada. Além disso, foi implementada uma estratégia de memoização para armazenar em cache os resultados de satisfatibilidade de fórmulas já processadas. A combinação dessas duas melhorias reduziu drasticamente o tempo de computação.

Para validar este avanço, realizamos um benchmark comparando o tempo de geração de três SAT solvers em 50 execuções, cada uma gerando 50 exercícios de complexidade 4 com todas as restrições ativadas. Os resultados, apresentados no gráfico de violino da Figura 2, mostram que o Picosat não apenas possui uma mediana de tempo de execução inferior, mas também uma distribuição muito mais consistente e com menos outliers, confirmando sua superioridade para esta aplicação.

Como resultado direto da geração, o GREAT produz não apenas a conjectura, mas também a estrutura da prova ou da refutação, como ilustrado na Figura 4, que mostra a implementação da prova em Coq e a árvore de prova correspondente. Isso é de grande valor para os instrutores, que podem verificar rapidamente a solução.

Em suma, os resultados obtidos transformaram o GREAT de um protótipo funcional para uma ferramenta madura e pronta para uso em larga escala. A combinação de uma interface amigável, geração rápida e saídas personalizáveis o torna uma solução eficaz para os desafios do ensino de lógica em turmas grandes, promovendo um aprendizado mais ativo e honesto. A disponibilização do código-fonte em github.com/terrematte/great reforça o compromisso do projeto com a comunidade acadêmica.

Conclusão

Ao longo do último ano, o projeto GREAT alcançou seus principais objetivos, evoluindo de uma ferramenta de nicho para uma aplicação web de código aberto, robusta e pronta para ser amplamente adotada. As melhorias implementadas abordaram diretamente as necessidades de educadores, focando em usabilidade, performance e flexibilidade.

Os resultados mais expressivos foram a otimização do tempo de geração de exercícios, validada por benchmarks que demonstraram a superioridade do SAT solver Picosat, e a adição de funcionalidades de exportação em LaTeX, que automatizam a criação de materiais didáticos de alta qualidade e dificultam a desonestidade acadêmica. A nova interface web tornou a ferramenta acessível a um público mais amplo, que pode agora gerar listas de exercícios personalizadas com poucos cliques.

O GREAT se destaca de outras propostas de AQG por não se limitar a equivalências proposicionais, mas gerar tarefas complexas de prova com premissas e conclusão, em múltiplos formatos.

Como trabalhos futuros, planejamos expandir a customização para templates de outros assistentes de prova, como o Lean, e otimizar ainda mais o algoritmo de geração. Além disso, pretendemos incorporar a estratégia de geração de exercícios com complexidade similar, baseada no número de regras de inferência necessárias, permitindo um controle ainda mais fino sobre o nível de dificuldade das tarefas geradas. Com isso, esperamos consolidar o GREAT como uma ferramenta indispensável no ensino de lógica e matemática para a computação.

Referências

Referências

- [1] Kurdi, G., Leo, J., Parsia, B. et al. A Systematic Review of Automatic Question Generation for Educational Purposes. *Int J Artif Intell Educ* 30, 121–204 (2020).<https://doi.org/10.1007/s40593-019-00186-y>.
- [2] Manoharan, S. Personalized Assessment as a Means to Mitigate Plagiarism. in *IEEE Transactions on Education*, vol. 60, no. 2, pp. 112-119, May 2017,<https://doi.org/10.1109/TE.2016.2604210>.
- [3] Srijek, Kiky, et al. Academic misconduct: Evidence from online class. *Int J Eval & Res Educ* 11.4 (2022): 1893-1902.<http://doi.org/10.11591/ijere.v11i4.23556>.
- [4] Yang, Yicheng, et al. Automatic question generation for propositional logical equivalences. *arXiv preprint* (2024).<https://doi.org/10.48550/arXiv.2405.05513>.
- [5] Terrematte, P.; Marcos, J. TryLogic tutorial: an approach to Learning Logic by proving and refuting. *Proceedings of the Fourth International Conference on Tools for Teaching Logic* (2015).<https://doi.org/10.48550/arXiv.1507.03685>.
- [6] Mendes, João; Marcos, João. A method for automated generation of exercises with similar level of complexity. In: *Workshop Brasileiro de Lógica (WBL)*. João Pessoa/PB. Anais. Porto Alegre: Sociedade Brasileira de Computação, p. 17-24 (2023).<https://doi.org/10.5753/wbl.2023.230660>.
- [7] Biere, A. PicoSAT Essentials. In: *Journal on Satisfiability, Boolean Modeling and Computation* 4 (2008), pp. 75-97.

Anexos

Box Plot - Figura 1

Web Interface

Exemplo - Proof Tree

CÓDIGO: HS0999

AUTOR: DÊNIS ROCHA DA SILVA

ORIENTADOR: DENNYS LEITE MAIA

TÍTULO: Criação de Objetos de Aprendizagem para Matemática com Inteligência Artificial Generativa

Resumo

O presente trabalho contempla o desenvolvimento de Objetos de Aprendizagem (OA) nas formações da segunda fase do projeto Inovática (Inovação Educacional em Matemática), promovidas nas unidades presenciais do Sistema Universidade Aberta do Brasil (UAB) no Rio Grande do Norte. O projeto, financiado pela Coordenação de Aperfeiçoamento de Pessoal de Nível Superior (CAPES), teve como objetivo capacitar professores da Educação Básica e estudantes de licenciatura para o uso de tecnologias digitais em sala de aula, explorando recursos inovadores como a Inteligência Artificial Generativa (IAG). A metodologia incluiu a apresentação da Plataforma OBAMA, a discussão sobre OAs e ensino de matemática, a exploração de Objetos de Aprendizagem (OAs) existentes e, por fim, a criação de OAs utilizando ferramentas de IAG. Foram produzidos um total de 31 OAs, 27 deles ligados à Matemática e foram classificados segundo as unidades temáticas da Base Nacional Comum Curricular (BNCC). Em síntese, as formações contribuíram para acrescentar mais um recurso às práticas docentes dos participantes, estimulando o protagonismo deles, enquanto educadores, na busca por construir estratégias pedagógicas que atendam às suas necessidades específicas.

Palavras-chave: Objeto de Aprendizagem, Matemática, Inteligência Artificial

TITLE: Creation of Learning Objects for Mathematics with Generative Artificial Intelligence

Abstract

The present work encompasses the development of Learning Objects (LO) during the training sessions of the second phase of the Inovática project (Educational Innovation in Mathematics), carried out at the in-person units of the Open University of Brazil (UAB) in Rio Grande do Norte. The project, funded by the Coordination for the Improvement of Higher Education Personnel (CAPES), aimed to train Basic Education teachers and undergraduate students in the use of digital technologies in the classroom, exploring innovative resources such as Generative Artificial Intelligence (GAI). The methodology included the presentation of the OBAMA Platform, discussions on LOs and mathematics teaching, the exploration of existing Learning Objects (LOs), and, finally, the creation of LOs using GAI tools. A total of 31 LOs were produced, 27 of them related to Mathematics, and

they were classified according to the thematic units of the Brazilian National Common Curricular Base (BNCC). In summary, the training sessions contributed to adding another resource to the teaching practices of the participants, fostering their protagonism as educators in the pursuit of building pedagogical strategies that meet their specific needs.

Keywords: Learning Objects, Mathematics, Artificial Intelligence

Introdução

O ensino de matemática é um grande desafio para educadores, em especial na atualidade, em que há uma vasta oferta de entretenimento disponível através de telas. Surge então a pergunta: como se pode tornar um conhecimento desenvolvido há milhares de anos atrativo para a nova geração? A resposta pode estar no uso da tecnologia para dar uma nova perspectiva à aprendizagem. Para tal, é importante que os professores possam se apropriar do uso dessas ferramentas tecnológicas e também disseminá-las para seus colegas.

Nesse escopo se encaixa o Projeto Inovática, que faz parte da iniciativa “Disseminação do Ensino de Matemática por meio da Plataforma OBAMA: Inovação, Formação e Colaboração”, financiado pela Coordenação de Aperfeiçoamento de Pessoal de Nível Superior - CAPES (Edital Inova EAD/CAPES nº 15/2023), que buscou disseminar práticas pedagógicas inovadoras através da Plataforma Objetos de Aprendizagem para a Matemática (OBAMA), a qual já foi bem descrita por Batista et al. (2017), Oliveira et al. (2018), Bezerra et al. (2019). Esse projeto foi dividido em duas etapas, englobando estudantes de licenciaturas da educação a distância da Universidade Federal do Rio Grande do Norte (UFRN) e professores da rede de ensino básico de municípios próximos aos polos da Universidade Aberta do Brasil (UAB).

A Plataforma OBAMA é um referatório virtual para objetos de aprendizagem, focados em matemática, lançado em maio de 2017 (Batista et al., 2017) e se propõe a ser um sítio de busca em que os professores consultem OAs e os utilizem em sala de aula. Os objetos disponibilizados na plataforma já estão classificados segundo as habilidades da BNCC e os descritores dos PCN (que é uma classificação legado), o que facilita a experiência de busca dos usuários, direcionando eles para os objetos que mais se adequam às suas demandas. Entretanto, em virtude da perda de OAs referenciados, tanto pela indisponibilidade dos recursos, quanto por obsolescência tecnológica, como, por exemplo, relatado por Farias et al. (2021) para o caso da tecnologia Flash, optou-se por explorar com os formandos do Inovática a criação de novos OAs, visando atender a necessidades específicas, além de incluí-los no acervo da OBAMA.

Recentemente, uma área que vem recebendo destaque é a Inteligência Artificial (IA), em especial, uma de suas novas vertentes, a chamada Inteligência Artificial Generativa (IAG), muito usada para criação de textos, imagens, vídeos e áudios. Entretanto, essa área é bem mais antiga, tendo o termo IA sido cunhado na década de 1950 em estudos de redes neurais (Russell & Norvig, 2013), influenciados pela ideia de que máquinas conseguem pensar sozinhas, discussão que é bem explorada por Turing (1950), claro, baseado na tecnologia que dispunha na época.

Entre as ferramentas de IAG disponíveis, é típico a necessidade de inserir um comando de entrada (prompt), geralmente um texto ou uma imagem, e é gerada uma resposta compatível

com a descrição fornecida. Entretanto, em tarefas mais elaboradas, nota-se a necessidade de detalhar as mídias de entrada, de maneira a evitar respostas inadequadas. Castilho, Rodriguez & Herrera (2024) sugerem uma classificação em três níveis para a relevância das informações contidas no que eles chamam de “prompt bem desenvolvido”: (1) necessário - Tarefa/Objetivo; (2) importante - Contexto/Detalhes, Exemplos; e (3) opcional - Público-Alvo, Formato, Persona e Diálogo Ativo.

Este trabalho visa apresentar os OAs gerados durante as formações presenciais, da segunda fase do Inovática, que se propuseram a auxiliar estudantes da área de educação e docentes do Rio Grande do Norte a criar comandos textuais de entrada para IAG destinados à produção de código de OAs, propiciando que eles criem objetos inéditos, ou adaptem alguns que já existem, ou ainda procurem atender necessidades específicas de acessibilidade. Além disso, esses objetos foram classificados segundo as unidades temáticas da BNCC, já pensando em inseri-los na OBAMA.

Metodologia

No geral, as formações ocorreram em oito polos da UAB que ofertam cursos de Licenciatura ou Pedagogia e estavam organizadas em quatro momentos, o primeiro compreendeu: Apresentação da plataforma e introdução ao conceito de OA, posteriormente, o uso de OA para o ensino de matemática, além da relevância do uso de metodologias ativas em sala de aula com auxílio de tecnologias digitais. Esses momentos foram importantes para que os participantes compreendessem os conceitos, além de propiciar a reflexão quanto ao uso desses recursos no cotidiano escolar. O segundo momento, com teor mais prático, englobava a exploração das funcionalidades da plataforma, reproduzindo como os formandos poderiam usá-la para consultar um OA que abrange um determinado tema. Com esse intuito, os formandos foram convidados a pensar em algum tema da matemática, um nível de ensino, ou mesmo em alguma habilidade da BNCC e acessar o site da plataforma em busca do OA que mais se aproximasse do que eles queriam. Em seguida, eles deveriam avaliá-los criticamente quanto a facilidade de interação, levantar formas de uso em aula, e relatar se ele cumpre o que propõe, ou ainda se enxergavam outros temas da matemática que o objeto também abrange.

A terceira parte, foi focada em IAG, quanto ao seu funcionamento e apresentando como opera na prática, além de reforçar a necessidade de detalhamento na elaboração de um prompt eficaz e com maior chance de retornar resultados satisfatórios, de forma a atender os níveis sugeridos por Castilho, Rodriguez & Herrera (2024), exemplificando: (1) Requisitar a criação de código em HTML, Javascript e CSS que retorne a interface de um OA para estudantes de algum nível do ensino básico; (2) Mostrar um tema ou conteúdo abordado, junto com a indicação da habilidade da BNCC ou tema de estudo no ano escolar a qual o OA se encaixa; e (3) Recomendar situações em que é apropriado o uso do OA em sala, levando em conta o público-alvo, além de demais informações competentes para o uso do OA, como, por exemplo, mídias que precisam estar presentes.

Um exemplo de prompt criado foi:

"Atue como um programador experiente, gere o código em html, css e javascript, em um arquivo só, de um objeto de aprendizagem de frações para estudantes do 4º ano do Ensino Fundamental, onde o usuário ligue a fração com a figura correspondente, é importante que

exiba uma mensagem de sucesso se o usuário acertar ou de falha se o usuário errar, além disso desejo que tenha um sistema de níveis, que vá aumentando a dificuldade sempre que o usuário acertar, podem ser três níveis no total, também é bom ter um botão para reiniciar quando chegar ao final. A figura que representa a fração pode ser em formato de pizza, estando pintadas apenas as partes equivalentes a fração, por exemplo, se a fração for $\frac{2}{3}$, aparece 2 pedaços da pizza pintadas e o restante fica branco".

Na quarta e última etapa da oficina, foi realizada a atividade final: Propor aos formandos a criação de um OA, baseado nos processos que foram apresentados nas etapas anteriores. Com tal intuito, eles foram estimulados a pensar como o objeto seria e a que se destinaria. Posteriormente, deveriam converter essa ideia em um comando de entrada e usá-lo no ChatGPT, ajustando as instruções até que o planejamento estivesse bem estruturado, e assim pedir que fosse gerado um código em HTML, Javascript e CSS em um arquivo só, de maneira a facilitar a execução através de ferramentas de edição online (como o W3Schools - https://www.w3schools.com/html/tryit.asp?filename=tryhtml_editor). Após isso, teve início a fase de testagem, que era adotada até que todos os erros verificados fossem solucionados.

A seção seguinte contempla a análise preliminar dos objetos criados nas formações. Esses recursos educacionais, em comum acordo com os autores, serão submetidos a melhorias e disponibilizados na Plataforma OBAMA.

Resultados e Discussões

Uma vez que o código era executado, os participantes puderam interagir com a interface criada, de forma a verificar o seu funcionamento e identificar comportamentos inadequados, assim, esses erros poderiam ser relatados ao ChatGPT, em uma nova instrução de comando com sugestões para a correção, o que propiciou aos formandos a experiência do Construcionismo (Papert, 2008), ao avaliarem as formas em que suas propostas de resolução de problemas eram compreendidas pelo computador, promovendo ainda as ações segundo o ciclo: descrição-execução-reflexão-depuração-abstração reflexionante (Valente, 1999). Cabe ressaltar, que essa etapa de testagem permitia aos formandos utilizarem as suas expertises, enquanto educadores, para atestar que os conhecimentos implementados pela IA eram verídicos ou que estavam sendo operados adequadamente, logo, por mais que o código estivesse sendo desenvolvido por IA, foi indispensável o papel do ser humano por trás da criação do prompt para reforçar positivamente às respostas que eram retornadas e garantir o direcionamento rumo ao resultado esperado.

Uma fragilidade observada foi quanto a uso de mídias como imagens e áudios, pois muitas vezes os links que o código consultava não estavam mais disponíveis, ou, no caso de imagens que eram geradas como gráficos através do Javascript, notava-se alguma inadequação com o esperado. Nesses casos, a recomendação foi buscar imagens mais adequadas na internet e fornecer o link delas, além de informar onde ela deveria ser inserida na próxima instrução textual. Já no caso de áudios que foram propostos por alguns dos formandos, esse processo de consulta não é tão simples, e o melhor caminho apresentado seria a gravação deles pelo próprio autor, o que demandaria muito tempo, logo eles optaram por não prosseguir com essa ideia.

Dentre os 31 OAs resultantes das oficinas, 26 deles tem algum tema da matemática abordado, e é essa a amostra que receberá enfoque neste trabalho. Com esse intuito, foi realizada uma categorização com base nas unidades temáticas da Base Nacional Comum Curricular

(BNCC) para a área de Matemática (BRASIL, 2018). As classificações de cada OA podem ser visualizadas na quarta coluna do Quadro 1. Nota-se uma predominância significativa da unidade temática "Números", presente em 20 dos OAs, o que representa mais da metade dos conteúdos abordados. Outras unidades que também se destacaram foram "Geometria" e "Probabilidade e Estatística", cada uma aparecendo em três OAs. Em menor frequência, o tema "Grandezas e Medidas" foi identificado apenas uma vez. Esses dados demonstram uma variedade de abordagens temáticas, embora com uma concentração mais

Outro tópico notável observado diz respeito às percepções dos participantes ao fim do processo de desenvolvimento dos OAs. Muitos relataram surpresa e entusiasmo ao perceberem o resultado final de seus trabalhos, especialmente por não imaginarem, inicialmente, que seriam capazes de criar recursos educacionais digitais e mesmo que usassem ferramentas de IA de outras formas, não haviam tido essa percepção de enxergá-las enquanto programadoras de objetos idealizados por eles. Além disso, compreenderam a importância de sua atuação nesse contexto, pois é quem escreve o prompt quem detém a capacidade de guiar a IA positivamente para a obtenção do melhor resultado possível.

Conclusão

O presente trabalho teve como propósito relatar e refletir sobre a experiência vivenciada no âmbito do projeto Inovática (Inovação Educacional em Matemática), com ênfase no processo formativo voltado à produção de Objetos de Aprendizagem (OAs). A análise concentra-se nos OAs que abordam temáticas da Matemática, que compõem a maioria dos produtos elaborados durante as formações. O estudo busca destacar as estratégias pedagógicas adotadas ao longo do percurso formativo e os principais resultados alcançados, evidenciando as contribuições para a prática docente e o papel da Inteligência Artificial Generativa (IAG) como recurso de apoio na criação dos OAs voltados a temas da Matemática.

Nas formações realizadas nos diferentes pólos do Sistema UAB no Rio Grande do Norte, a proposta de produção de OAs se apresentou como uma abordagem inovadora para o planejamento de práticas pedagógicas criativas e contextualizadas. A utilização da IAG nesse processo possibilitou aos participantes expandirem suas formas de pensar a abordagem de assuntos da Matemática, incorporando recursos digitais como ferramentas para auxiliar a transmissão de conteúdos que muitas vezes são de difícil assimilação para os alunos.

A utilização da IAG mostrou-se uma aliada estratégica na elaboração dos OAs de Matemática, pois, ao eliminar barreiras técnicas como programação ou design gráfico, a ferramenta permitiu que os educadores direcionassem seus esforços para a elaboração de propostas pedagógicas alinhadas às demandas reais de suas salas de aula. Esse aspecto tornou o processo mais acessível e reforçou a ideia de que professores podem desenvolver recursos didáticos digitais sem depender exclusivamente de especialistas técnicos.

Por conseguinte, conclui-se que iniciativas formativas como esta são essenciais para impulsionar a inovação nos processos educativos, e se espera que as reflexões e resultados positivos apresentados possam servir como inspiração para ampliar futuras ações formativas que incorporem as tecnologias educacionais de forma crítica, criativa e pedagógica, com ênfase no potencial transformador da IAG no ensino da Matemática.

Referências

Batista, S., Bezerra, V., Oliveira, A. & Maia, D. (2017). OBAMA: Um repositório de objetos de aprendizagem para matemática. In Anais dos Workshops do Congresso Brasileiro de Informática na Educação, 300. DOI: <https://doi.org/10.5753/cbie.webie.2017.300>.

Bezerra, V., Batista, S., Oliveira, A., & Maia, D. (2019). Processo de identificação e correção de erros da plataforma OBAMA. In Anais do Workshop de Informática na Escola (pp. 1384–1388). SBC. <https://sol.sbc.org.br/index.php/wie/article/download/13322/13175/>.

Brasil. Ministério da Educação. (2018). Base Nacional Comum Curricular (BNCC). http://basenacionalcomum.mec.gov.br/images/BNCC_EI_EF_110518_versaofinal_s ite.pdf.

Castilho, G., Rodriguez, C. & Herrera, V. (2024). Um relato de experiência de aplicação de engenharia de prompt no ensino superior em STEM. In Workshop em Estratégias Transformadoras e Inovação na Educação (WETIE) (pp. 69–78). SBC.

Farias, F., Carvalho, A., Rodrigues, R., de Oliveira, N., & Maia, D. (2021). Impacto da descontinuidade da tecnologia Flash na disponibilidade dos objetos de aprendizagem em repositórios educacionais. In Anais do VI Congresso sobre Tecnologias na Educação (Ctrl+E 2021). (pp. 90–99). SBC. <https://sol.sbc.org.br/index.php/ctrl/article/view/17553>

Oliveira, A., Batista, S., Nascimento, Í., Azevedo, D., Lima, R., Oliveira, N. & Maia, D. (2018). Processo de desenvolvimento de uma ferramenta destinada à elaboração de planos de aula de forma colaborativa. In Anais do III Congresso sobre Tecnologias na Educação (CTRL+E 2018). UFC. https://ceur-ws.org/Vol-2185/CtrlE_2018_paper_118.pdf.

Papert, S. (2008). A máquina das crianças: Repensando a escola na era da informática (S. Costa, Trad.). Artmed. (Trabalho original publicado em 1993)

Russell, S., & Norvig, P. (2013). Inteligência artificial (R. C. Simille, Trad.). Elsevier.

Timpone, R., & Guidi, M. (2023). Explorando a mudança de cenário da IA: Da IA analítica à IA generativa. Ipsos Knowledge Centre.

TURING, Alan M. Computing Machinery and Intelligence. Mind, New Series. Oxford University Press on behalf of the Mind Association, v. 59, n. 236, p. 433-460, 1950.

Valente, J. A. (1999). O computador na sociedade do conhecimento. UNICAMP/NIED.

Anexos

Figura 1. Formação no Polo UAB de Natal - RN.

Figura 2. Formação no Polo UAB de Martins - RN.

Figura 3. Ciclo do Construcionismo vivenciado durante a elaboração de OAs.

Figura 4. Interface do OA “Jogo da Memória Junino”.

Figura 5. Interface do OA “Estatística: Média, Mediana e Moda - Interativo”.

Figura 6. Gráfico que mostra os conteúdos mais abordados pelos OAs.

Quadro 1. OAs feitos na formação, nome dos autores e conteúdos explorados.

CÓDIGO: HS1110

AUTOR: SAMUEL ANDERSON MACHADO LOPES

ORIENTADOR: DENNYS LEITE MAIA

TÍTULO: A Missão de Rahim: Validação de um Objeto de Aprendizagem para o Ensino de Matemática

Resumo

O artigo discute a validação do Objeto de Aprendizagem Batalha de Algarismos: A Missão de Rahim, a partir da aplicação de um questionário a estudantes do 6º ano do Ensino Fundamental. O instrumento, elaborado em formato digital, contou com questões de múltipla escolha e abertas, permitindo avaliar dimensões como usabilidade, clareza das instruções, compreensão dos conteúdos e motivação dos alunos. A análise descritiva das respostas evidencia a confiabilidade e a validade do recurso, além de indicar ajustes necessários para seu aprimoramento. Os resultados evidenciam a validação do objeto de aprendizagem, destacando seu potencial pedagógico ao motivar os alunos pelo caráter lúdico.

Palavras-chave: Objeto de Aprendizagem; Validação; Ensino Fundamental; Usabilidade;

TITLE: Rahim's Mission: Validation of a Learning Object for Mathematics Teaching

Abstract

This article discusses the validation of the Learning Object "Battle of Numbers: Rahim's Mission," based on a questionnaire administered to sixth-grade students. The instrument, designed in digital format, included multiple-choice and open-ended questions, allowing for the assessment of dimensions such as usability, clarity of instructions, content comprehension, and student motivation. A descriptive analysis of the responses demonstrates the resource's reliability and validity, as well as indicating necessary adjustments for improvement. The results demonstrate the validation of the learning object, highlighting its pedagogical potential by motivating students through its playful nature.

Keywords: Learning Object; Validation; Elementary Education; Usability;

Introdução

Objetos de Aprendizagem (OAs) são recursos digitais disponibilizados na web, reutilizáveis e voltados para a aprendizagem de um conteúdo específico (Rebouças, Maia e Scaio, 2021). Suas características incluem facilidade de atualização, customização, interoperabilidade e tamanho reduzido. A utilização de objetos de aprendizagem no ensino se configura como uma estratégia capaz de articular competências técnicas e conhecimentos pedagógicos os quais possibilitam experiências educativas mais dinâmicas e significativas no ambiente educacional. Como destacam Maia et al. (2017), OAs podem contribuir para a diversificação das práticas pedagógicas ao permitir diferentes representações e possibilidades de

manipulação conceitual, reforçando a importância do design pedagógico no processo de criação.

De acordo com Oliveira, Costa e Moreira (2001), a avaliação de softwares educativos e ambientes informatizados de aprendizagem observa os aspectos de interação, fundamentação pedagógica, pertinência do conteúdo e estabilidade técnica, garantindo que o recurso digital não apenas funcione de maneira eficiente, mas também esteja alinhado a objetivos educacionais claros. A avaliação formativa, acompanhada da aplicação em contextos reais, permite compreender em que medida o objeto de aprendizagem favorece a aprendizagem significativa, além de identificar fragilidades que podem ser corrigidas em versões posteriores.

Nesse ínterim que se insere o presente artigo, que discute a validação do Objeto de Aprendizagem Batalha de Algarismos: A Missão de Rahim. Aqui busca-se compreender como o recurso educativo digital é avaliado pelos estudantes, em que medida contribui para o processo de ensino e aprendizagem e quais ajustes podem potencializar sua utilização em contextos escolares.

Metodologia

A motivação para produção do OA “Batalha de Algarismos: A Missão de Rahim” (Lopes et al., no prelo) decorreu da redução de objetos de aprendizagem para Matemática devido à descontinuidade da tecnologia Flash na Plataforma Obama (Farias et al., 2021). Nesse contexto, a equipe de desenvolvimento, a partir de atividades de um curso técnico em Tecnologia da Informação (TI) e da participação em grupo de pesquisa e desenvolvimento sobre tecnologias educacionais, propôs um OA voltado para a leitura e escrita de algarismos romanos. A relação entre a História e a Matemática encontra nos Algarismos romanos um espaço em que ambas se complementam. De acordo com Sousa (2023), a Aliança entre História da Matemática (HM) e Tecnologias Digitais (TD) é uma das tendências promissoras em Educação Matemática pois busca explorar o aspecto do desenvolvimento histórico da Matemática por meio de ferramentas e recursos digitais, visando uma aprendizagem mais dinâmica e envolvente. Nesse sentido, o OA desenvolvido oportuniza a manipulação de representações dos conceitos, diversificando as estratégias de trabalho de professores que ensinam Matemática.

O “Batalha de Algarismos” foi desenvolvido, entre novembro de 2024 e maio de 2025, utilizando HTML, CSS e JavaScript, ferramentas de código aberto que permitem carregamento leve e acesso tanto online quanto por download, sendo licenciado sob Creative Commons NY-BY-CC. Esta licença permite que os reutilizadores distribuam, remixem, adaptem e desenvolvam o material em qualquer meio ou formato apenas para fins não comerciais e somente enquanto a atribuição for dada ao criador. Para a produção das imagens, foram utilizadas ferramentas de inteligência artificial generativa (IA), como o ChatGPT, o Leonardo AI e o Copilot, calibradas conforme a intenção da equipe, garantindo que as imagens representassem a proposta histórico-romano-persa.

Como parte do processo de desenvolvimento e avaliação do OA, foi elaborado um questionário online, estruturado como formulário digital do Google, com perguntas predominantemente de múltipla escolha, complementadas por algumas questões abertas. O instrumento foi planejado para avaliar usabilidade, clareza das instruções, compreensão dos

conceitos matemáticos e motivação dos alunos, integrando a avaliação diretamente ao processo de criação do OA. Esse procedimento segue recomendações de Marconi e Lakatos (2003, p. 212) sobre a utilização de formulários para contato direto com os informantes, possibilitando coleta de informações detalhadas sobre a experiência dos alunos.

A população do estudo são alunos do 6º ano do Ensino Fundamental, público-alvo do OA, da Escola Estadual Floriano Cavalcanti, garantindo alinhamento entre o conteúdo do recurso e os objetivos de aprendizagem, conforme a Base Nacional Comum Curricular (BNCC). A amostra foi selecionada de forma conveniente, considerando o acesso facilitado aos respondentes e a possibilidade de aplicar o instrumento de pesquisa sem interferir na rotina escolar.

Os dados coletados foram analisados por meio de estatística descritiva básica, utilizando planilhas eletrônicas e os recursos do Google Forms para compilar, organizar e comparar respostas.

Resultados e Discussões

O questionário de avaliação do objeto de aprendizagem “Batalha de Algarismos” foi respondido por 17 estudantes do ensino fundamental, permitindo uma análise inicial do impacto do recurso digital sobre o aprendizado dos alunos. Há distribuição de gênero, 47,1% dos respondentes eram do sexo masculino e 52,9% do sexo feminino, demonstrando uma participação equilibrada que possibilita observar diferentes perspectivas de aprendizado e engajamento. Quanto à faixa etária predominante foi de 11 a 13 anos, representando 88,2% dos participantes, enquanto 11,8% tinham entre 14 e 17 anos. Como mostrado no Gráfico 1 e Gráfico 2 respectivamente.

Em relação à compreensão do conteúdo matemático, 94,1% dos alunos relataram que conseguiram compreender a representação de números do sistema decimal no sistema romano após a utilização do objeto. Esse resultado indica que a maioria dos estudantes conseguiu fazer a transposição de conceitos abstratos para um contexto interativo e lúdico, favorecendo a aprendizagem significativa. Apenas 5,9% dos alunos indicaram não ter compreendido a relação entre os sistemas numéricos, o que sugere a necessidade de ajustes pontuais no material ou na mediação pedagógica para atender a todos os níveis de aprendizado. Como está apresentado no Gráfico 3.

O interesse dos estudantes pelo tema também foi altamente positivo. Todos os participantes afirmaram que a interação com o recurso digital despertou motivação para aprender mais sobre o sistema de numeração romano. Esse dado evidencia que o objeto de aprendizagem não apenas oportuniza apropriação do conceito, mas também consegue engajar os alunos, tornando o processo educativo mais atrativo.

Quanto à estética, narrativa e ambientação do OA, todos os alunos consideraram os cenários, personagens e diálogos adequados ao contexto da proposta, o que reforça a importância do design gráfico e da narrativa na aprendizagem digital. A avaliação subjetiva sobre a aplicabilidade em sala de aula revelou que os estudantes desejam que práticas com recursos educacionais digitais (REDs) sejam incorporadas pelos professores, indicando uma percepção positiva sobre a integração da tecnologia no processo educativo.

Em relação à retratação histórica apresentada pelo OA, os resultados mostraram que 58,8% dos estudantes não identificaram ou não aprenderam sobre os aspectos históricos incluídos, enquanto 41,2% reconheceram a representação histórica presente no material. Esse dado sugere que, embora o foco principal do recurso digital seja o aprendizado matemático, há espaço para aprimorar a contextualização histórica de forma mais explícita e integrada às atividades. Como evidência o Gráfico 4.

A dinâmica do Recurso educacional foi amplamente avaliada de forma positiva, sendo considerada atrativa e motivadora para a aprendizagem. No entanto, os estudantes também apresentaram sugestões para melhorias técnicas, como aumento do nível de dificuldade, maior tempo disponível para realização das fases, inclusão de mais fases e outras recomendações. Essas observações fornecem insumos importantes para ajustes futuros, garantindo que o OA continue estimulante e adequado a diferentes perfis de estudantes. Vale salientar que, ao longo do processo de avaliação, os alunos reconheceram o OA predominantemente como um jogo, e não apenas como um recurso educacional. Essa percepção, não representou uma limitação ao potencial do recurso educacional, uma vez que a familiaridade e o hábito com jogos digitais contribuíram para aumentar a motivação dos estudantes em experimentar o recurso educacional. Esse caráter lúdico favoreceu o engajamento, facilitou a aprendizagem e despertou maior interesse pelo conteúdo, ao aproximar o processo educativo de práticas já presentes no cotidiano dos alunos.

Conclusão

A validação realizada com estudantes do ensino fundamental indicou que o OA Batalha dos Algarismos contribui significativamente para o desenvolvimento de habilidades matemáticas e do raciocínio lógico, ao mesmo tempo em que desperta interesse pelo tema por meio de sua.

Os resultados da avaliação evidenciam que a maioria dos alunos compreendeu a representação de números do sistema decimal no sistema romano e considerou os cenários, personagens e diálogos adequados ao contexto da proposta. A dinâmica do OA foi avaliada como atrativa e motivadora para a aprendizagem, reforçando a importância da integração entre diferentes áreas do conhecimento na criação de recursos educacionais eficazes. Além disso, os estudantes demonstraram interesse em utilizar práticas com recursos educacionais digitais em sala de aula, evidenciando o potencial do OA como ferramenta pedagógica.

Apesar dos resultados positivos, a avaliação também apontou pontos que podem ser aprimorados. Entre as sugestões, destacam-se o aumento do nível de dificuldade das fases, a ampliação do tempo disponível para a realização das atividades, a inclusão de mais fases e ajustes técnicos gerais para tornar a experiência ainda mais desafiadora e envolvente, também foi identificada a necessidade de ajustes nos elementos visuais, de modo a potencializar a clareza, a atratividade e a adequação estética do recurso. Essas recomendações fornecem subsídios importantes para o refinamento do OA, garantindo que ele continue estimulante e eficaz para diferentes perfis de estudantes.

Referências

CASTRO-FILHO, J.; CASTRO, J.; SOUZA, M.; FREIRE, R.; NASCIMENTO, G. O reino de Aljabar: o desafio da balança: um recurso educacional digital para favorecer o

desenvolvimento do pensamento algébrico. In: CONGRESSO BRASILEIRO DE INFORMÁTICA NA EDUCAÇÃO, 10., 2020, [S. l.]. Anais [...]. Porto Alegre: SBC, 2020. p. 197-204. DOI: <https://doi.org/10.5753/wcbie.2021.218536>

FARIAS, F.; CARVALHO, A.; RODRIGUES, R.; OLIVEIRA, N.; MAIA, D. Impacto da descontinuidade da tecnologia Flash na disponibilidade dos objetos de aprendizagem em repositórios educacionais. In: CONGRESSO SOBRE TECNOLOGIAS NA EDUCAÇÃO, 6., 2021, [S. l.]. Anais [...]. Porto Alegre: SBC, 2021. p. 90–99. Disponível em: <https://sol.sbc.org.br/index.php/ctrl/article/view/17553>. Acesso em: 30 ago. 2025.

MAIA, D.; OLIVEIRA, A.; SILVA, A.; COSTA, C.; BRITO, D.; MELO, E.; OLIVEIRA, N. Objetos de aprendizagem para matemática: Yes we can! In: CONGRESSO SOBRE TECNOLOGIAS NA EDUCAÇÃO, 2., 2017, João Pessoa. Anais [...]. João Pessoa: UFPB, 2017. Disponível em: https://ceur-ws.org/Vol-1877/CtrlE2017_MC_10.pdf. Acesso em: 30 ago. 2025.

MARCONI, M. A.; LAKATOS, E. M. Fundamentos de metodologia científica. 7. ed. São Paulo: Atlas, 1992.

OLIVEIRA, A.; BARRETO, G.; SANTOS, M.; SILVA, L.; SILVA, I.; SOUSA, G.; MAIA, D. O desafio de Al-Biruni: solucionando o problema do xadrez. In: CONGRESSO BRASILEIRO DE INFORMÁTICA NA EDUCAÇÃO, 2024, [S. l.]. Anais Estendidos [...]. Porto Alegre: SBC, 2024. Disponível em: https://sol.sbc.org.br/index.php/cbie_estendido/article/download/31777/31579/. Acesso em: 30 ago. 2025.

OLIVEIRA, A.; MARINHEIRO, F.; CUNHA, I.; PEREIRA, M.; KIMURA, R.; MAIA, D. Software educativo Brincáculo. In: CONGRESSO REGIONAL SOBRE TECNOLOGIAS NA EDUCAÇÃO, 2016, Natal. Anais [...]. Natal: UFRN, 2016. p. 640-645. Disponível em: https://ceur-ws.org/Vol-1667/CtrlE_2016_MS_paper_72.pdf. Acesso em: 30 ago. 2025.

OLIVEIRA, C.; COSTA, J.; MOREIRA, M. Ambientes informatizados de aprendizagem: produção e avaliação de software educativo. Campinas: Papirus, 2001.

PIRES, A.; COLAÇO, H.; HORTA, M.; RIBEIRO, C. Desenvolver o sentido de número no pré-escolar. Faro: UAIG, 2013. Disponível em: <https://core.ac.uk/download/pdf/61518672.pdf>. Acesso em: 30 ago. 2025.

SANTOS, M.; SOUSA, G. Aliança entre história da matemática e tecnologias digitais numa proposta de produto educacional baseado na obra A cronologia das nações antigas de Al-Biruni. In: CONGRESSO NACIONAL DE EDUCAÇÃO, 2024, Campina Grande. Anais [...]. Campina Grande: Realize Editora, 2024. v. 2. Disponível em: <https://editorarealize.com.br/artigo/visualizar/105746>. Acesso em: 30 ago. 2025.

SOUSA, G. História da matemática em alianças com tecnologias digitais. REMATEC, Natal, v. 18, n. 44, p. 1-18, 2023. DOI: <https://doi.org/10.37084/REMATEC.1980-3141.2023.n44.pe2023005.id510>

Anexos

Gráfico 1: Distribuição de gênero

Gráfico 2: Faixa etária

Gráfico 3: Compreensão do conteúdo matemático

Gráfico 4: Compreensão do conteúdo histórica

CÓDIGO: HS1283

AUTOR: ANA CAROLINA COSTA SILVA

ORIENTADOR: DENNYS LEITE MAIA

TÍTULO: PRÁTICAS INVESTIGATIVAS E CRIATIVAS NA EDUCAÇÃO BÁSICA: UM OLHAR SOBRE A FORMAÇÃO DOCENTE NA REM-NE

Resumo

Este trabalho analisa como uma professora da rede pública de ensino percebe e reflete sobre práticas investigativas e criativas em sua atuação docente. O trabalho é fruto do acompanhamento das formações realizadas pela REM-NE, que tiveram como objetivo promover experiências formativas que contribuíssem para o desenvolvimento profissional docente. A metodologia consiste na análise de uma entrevista semiestruturada como principal instrumento de coleta de dados, com o intuito de identificar as percepções individuais da professora e compreender como a participação nas formações contribuiu para o aprimoramento da sua atuação em sala de aula. Por fim, o estudo destacou que experiências formativas podem fortalecer e transformar a prática docente.

Palavras-chave: Formação docente. Planejamento. Abordagem.

TITLE: INVESTIGATIVE AND CREATIVE PRACTICES IN BASIC EDUCATION: A
LOOK AT TEACHER TRAINING IN REM-NE

Abstract

This study analyzes how a public school teacher perceives and reflects on investigative and creative practices in her teaching. The study is the result of monitoring training programs conducted by REM-NE, which aimed to promote formative experiences that contributed to the professional development of teachers. The methodology used consisted of analyzing a semi-structured interview as the primary data collection tool, aiming to identify the teacher's individual perceptions and understand how participation in the training contributed to improving her classroom performance. Finally, the study highlighted that formative experiences can strengthen and transform teaching practice.

Keywords: Teacher training. Planning. Approach.

Introdução

A formação continuada de professores é condição essencial para o processo constante de aperfeiçoamento dos conhecimentos e competências essenciais à prática docente. De acordo com Libâneo (2014), diante das transformações sociais, culturais e tecnológicas, o trabalho do professor não deve restringir-se à mera transmissão de conteúdos; exige, ao contrário, um esforço contínuo de estudo, reflexão e atualização, capazes de promover condições para uma ação docente alinhada às necessidades dos seus estudantes.

Nesse contexto, a Rede de Educação Matemática do Nordeste (REM-NE), cujo tema orientador é Desenvolvimento Profissional do Professor em Ações Interdisciplinares com Equidade, configura-se como um importante espaço de formação e diálogo entre docentes e pesquisadores. O grupo promove encontros formativos que buscam fortalecer a prática pedagógica, estimular a reflexão crítica sobre as metodologias de ensino e favorecer a implementação de abordagens criativas e investigativas.

A participação de professores em espaços formativos favorece o fortalecimento de sua identidade profissional, uma vez que permite que o professor se reconheça como sujeito ativo de sua formação. Além disso, ao assumir uma postura investigativa, o docente desenvolve práticas pedagógicas autorais e contextualizadas, tornando-se protagonista na elaboração de experiências de aprendizagem que valorizam a construção coletiva do conhecimento.

Considerando esse cenário, levanta-se a seguinte questão: como uma professora da rede pública percebe e reflete sobre práticas investigativas e criativas em sua atuação docente? Busca-se compreender em que medida a participação nas formações da REM-NE contribui para o fortalecimento da identidade profissional e para o aprimoramento das estratégias de ensino, evidenciando a importância da formação continuada como espaço de transformação pedagógica e valorização docente.

Metodologia

Este estudo fundamenta-se como um estudo de caso de abordagem qualitativa, uma vez que busca aprofundar a compreensão de um fenômeno específico dentro de seu contexto real, considerando as experiências, percepções e práticas de uma professora da rede pública de ensino. De acordo com Stake (1994, apud ANDRÉ, 2013, p. 96), “estudo de caso não é uma escolha metodológica, mas uma escolha do objeto a ser estudado”, destacando que o conhecimento produzido é mais concreto, contextualizado e interpretativo, permitindo compreender o fenômeno em suas especificidades e valorizando a interpretação da participante sobre suas próprias experiências.

Nesse sentido, essa perspectiva evidencia que a escolha pelo estudo de caso se justifica pela necessidade de analisar em profundidade uma situação singular, considerando os significados construídos pela docente no contexto das formações promovidas pela REM-NE.

A coleta de dados ocorreu principalmente por meio de entrevista semiestruturada com a professora participante, permitindo capturar suas percepções e experiências sobre práticas investigativas e criativas de ensino. Também foram considerados os registros de anotações em diário de campo das pesquisadoras durante os encontros formativos.

O processo formativo investigado ocorreu entre junho e novembro de 2024, em uma escola pública de Parnamirim (RN), envolvendo momentos de discussão, planejamento de sequências de ensino e atividades que integraram teoria e prática, com foco na interdisciplinaridade e na inovação pedagógica. Esse acompanhamento permitiu observar de que forma a formação contribuiu para o fortalecimento da identidade profissional da docente e para a implementação de estratégias pedagógicas mais criativas e significativas em sala de aula.

Resultados e Discussões

Nesta seção, serão detalhados os resultados obtidos a partir da entrevista com a professora, buscando responder a questão central da pesquisa: como uma professora da rede pública de ensino percebe e reflete sobre práticas investigativas e criativas de ensino em sua atuação docente? A análise tem foco, sobretudo, como a participação nas formações promovidas pela REM-NE contribuiu para o refinamento das estratégias de ensino, destacando o impacto da formação continuada na prática pedagógica da professora.

Ao refletir sobre suas práticas, a professora evidenciou que a participação nas formações proporcionou a aplicação de práticas investigativas e criativas, especialmente na interdisciplinaridade entre as disciplinas de Matemática e Ciências. Ela aponta que “a interdisciplinaridade favoreceu várias coisas no aprendizado do aluno”, e que com as atividades práticas, os estudantes se mostraram motivados a ir além do que a professora esperava. (Entrevista, 1:28 - 1:36). Assim, essa percepção indica que a docente entende a prática investigativa não apenas como um momento de transmissão do conhecimento, mas como uma oportunidade do aluno experimentar, vivenciar e atribuir significado ao conteúdo que está sendo estudado. De acordo com Freire (1996), a aprendizagem significativa ocorre quando o professor estabelece uma relação entre a teoria e a prática, conectando os conteúdos ensinados à realidade dos alunos.

Além disso, a docente destacou que a abordagem possibilitou um cenário em que os estudantes se tornassem protagonistas no processo de aprendizagem, uma vez que “a partir do momento que ele está construindo, ele está aprendendo. Se todo mundo está construindo e tentando, é a forma mais simples de aprender” (Entrevista, 11:02 - 11:08).

Nesse contexto, a relevância da sequência de ensino planejada em conjunto com a professora se evidencia, especialmente na articulação interdisciplinar dos conteúdos, pois segundo a docente “A sequência de ensino não desassociou de nada, casou com tudo que a gente estava dando no momento” (Entrevista, 3:42 - 3:50). Essa fala demonstrou que o processo formativo elaborado pelo grupo da REM-NE não se limita à mera transmissão de conteúdos, mas apoia o professor na construção de estratégias pedagógicas contextualizadas e significativas para a prática em sala de aula.

Além disso, foi possível notar que o processo formativo contribuiu para o fortalecimento da sua identidade docente, uma vez que ao refletir sobre suas práticas e integrar uma novas estratégias de ensino de maneira intencional, a docente passa a se ver como alguém capaz de planejar, tomar decisões e adaptar atividades de acordo com a realidade dos seus estudantes. Como ela afirmou no trecho “Tudo que a gente fez na formação pode ser utilizado em qualquer momento na sala de aula. Porém, tem que ser planejado para ser executado” (Entrevista, 7:47 - 8:01).

Dessa forma, as formações promovidas pelo grupo da REM-NE mostra-se relevante não apenas pelo conteúdo abordado, mas pelo estímulo à reflexão, à inovação pedagógica e ao reconhecimento do professor como agente ativo na construção do aprendizados dos estudantes. Além disso, esse modelo formativo pode ser fonte de inspiração para outras iniciativas de formação continuada, pois demonstrou como é possível apoiar efetivamente o desenvolvimento docente e impactar sua prática em sala de aula.

Conclusão

O presente estudo revelou a importância da formação continuada no refinamento da prática pedagógica dos professores, especialmente quando oferece espaços de reflexão, inovação e diálogo entre os participantes. Dessa forma, a participação da professora nas formações da REM-NE proporcionou um novo olhar sobre sua prática docente, motivando-a a se tornar uma agente ativa na construção de experiências de aprendizagem mais significativas.

Observou-se que a incorporação de práticas investigativas e criativas, estimuladas durante as formações, manifestaram-se em sua atuação de forma concreta, seja no planejamento de atividades alinhadas ao contexto dos estudantes ou na valorização do protagonismo estudantil. Assim, a experiência vivenciada pela professora, reforça que a formação docente não deve limitar-se à transmissão de conteúdos, mas sim promover o desenvolvimento de práticas que impactam o fazer pedagógico dos professores.

Além disso, ficou evidente que a formação promovida pela REM-NE favorece a ressignificação da prática docente, uma vez que possibilitou a professora revisitar suas concepções e experimentar novas abordagens de ensino. Dessa forma, foi possível integrar reflexão e ação, aproximando teoria e prática.

Por fim, compreende-se que investir em formações semelhantes a REM-NE é fundamental para a valorização e fortalecimento da prática docente. Ao oportunizar o diálogo, a experimentação e a construção coletiva de saberes, essas formações ampliam as possibilidades de ensino e contribuem para a construção de uma educação mais crítica, criativa e alinhada às necessidades dos estudantes.

Referências

- LIBÂNEO, José Carlos. Adeus professor, adeus professora?. Cortez editora, 2014. Disponível em: https://books.google.com/books?hl=pt-BR&lr=&id=BOK_AwAAQBAJ&oi=fnd&pg=PT6&dq=LIB%C3%82NEO,+Jos%C3%A9+Carlos.+Adeus+professor,+adeus+professora%3F.+Cortez+editora,+2014.&ots=Dun-tb1KEx&sig=aT5C61p4dl8CXAHIg7VvAZPhEl0 . Acesso em: 20 de agos. de 2025.
- ANDRÉ, Marli. O que é um estudo de caso qualitativo em educação. Revista da FAAEBA: Educação e Contemporaneidade, p. 95-103, 2013. Disponível em: http://educa.fcc.org.br/scielo.php?pid=S0104-70432013000200009&script=sci_abstract&tlng=en . Acesso em: 20 de agos. de 2025.
- FREIRE, Paulo. Pedagogia da autonomia: saberes necessários à prática educativa. Editora Paz e terra, 2014. Disponível em: <https://books.google.com/books?hl=pt-BR&lr=&id=Ae4nAwAAQBAJ&oi=fnd&pg=PT3&dq=FREIRE,+Paulo.+Pedagogia+da+autonomia:+saberes+necess%C3%A1rios+%C3%A0+pr%C3%A1tica+educativa.+Editora+Paz+e+terra,+2014.&ots=MYcy5BXtcn&sig=ry4Qlc-5Jb0rnEuzXb-0RnV6IZo> . Acesso em: 20 de agos. de 2025.
- MAIA, Dennys Leite. Práticas pedagógicas inovadoras com diferentes tecnologias: conceitos e relatos de experiências. In: SILVA, Cristiane; MELLO, Roger. (Orgs.). Metodologias,

práticas e inovação na educação contemporânea. Volume 1. Rio de Janeiro: e-Publicar, 2021, p.312-331. Disponível em:
<https://www.dropbox.com/scl/fi/5dxss8fiufnd3c7qmq78d/Maia-2021-Pr-ticas-pedag-gicas-inovadoras-com-diferentes-tecnologias.pdf?rlkey=oh44poo39fwe5bc9ygyq8vnwt&e=1&dl=0>
. Acesso em: 20 de agos. de 2025.

CÓDIGO: HS1523

AUTOR: KEVEN WILLIAM PEREIRA MONTEIRO

ORIENTADOR: DENNYS LEITE MAIA

TÍTULO: APRIMORAMENTO DA BUSCA AVANÇADA DA PLATAFORMA OBAMA

Resumo

Este artigo trata do aprimoramento da funcionalidade de Busca Avançada da Plataforma OBAMA (Objetos de Aprendizagem para a Matemática), um referatório digital essencial para professores da Educação Básica. O trabalho é fruto de um projeto de Iniciação Científica e tem como objetivo reorganizar e transformar essa funcionalidade em um serviço mais eficiente e intuitivo, alinhado às necessidades do planejamento docente. Realizou-se o redesenho da interface do usuário, agrupando os filtros pedagógicos em uma janela modal, um componente de interface que conforme os princípios de design de interfaces expostos por Cybis et al. (2007) promove foco e eficiência sem perder o contexto principal, desenvolvida com a linguagem Dart e o framework Flutter. O estudo oportunizou constatar que a nova interface torna o processo de encontrar Objetos de Aprendizagem (OAs) mais rápido e preciso, fortalecendo a plataforma como uma ferramenta indispensável para a inovação na educação matemática.

Palavras-chave: Busca Avançada. Plataforma OBAMA. Experiência do Usuário (UX).

TITLE: Development of an Advanced Search Web Service for the OBAMA Platform

Abstract

This article addresses the enhancement of the Advanced Search feature for the OBAMA Platform (Learning Objects for Mathematics), an essential digital referatory for Basic Education teachers. This work is the result of a Scientific Initiation project and aims to reorganize and transform this functionality into a more efficient and intuitive service, aligned with the needs of teaching planning. The methodology involved redesigning the user interface by grouping pedagogical filters into a modal window, an interface component that according to the design principles outlined by Cybis et al. (2007) promotes focus and efficiency without losing the main context, developed with the Dart language and the Flutter framework. The study found that the new interface makes the process of finding Learning Objects (LOs) faster and more precise, strengthening the platform as an indispensable tool for innovation in mathematics education.

Keywords: Advanced Search. OBAMA Platform. User Experience (UX).

Introdução

A Plataforma Objetos de Aprendizagem para a Matemática (OBAMA) consolida-se como um referatório digital essencial para professores da Educação Básica, oferecendo um acervo curado de recursos educacionais digitais (REDs). Conforme destacado por Batista et al. (2017), seu principal objetivo é apoiar a diversificação de estratégias pedagógicas, superando as limitações dos buscadores tradicionais da internet, que retornam resultados nem sempre adequados ao contexto educacional.

No cerne dessa missão está a eficiência do mecanismo de descoberta de recursos. Um acervo robusto, que conta com centenas de OAs (Bezerra et al., 2019), é de pouca utilidade se os professores não conseguirem encontrar, de forma rápida e precisa, o recurso ideal para uma habilidade específica da BNCC ou um conteúdo que pretendem trabalhar. A funcionalidade de Busca Avançada emerge, portanto, não como um acessório, mas como um serviço estratégico e o coração da experiência do usuário na plataforma.

Este trabalho, desenvolvido no âmbito do projeto de Iniciação Científica, relata o redesenho e a reorganização desta funcionalidade crítica. O foco foi transformar a Busca Avançada em uma ferramenta verdadeiramente eficaz, organizando os filtros pedagógicos existentes de forma lógica e intuitiva, dentro de uma interface moderna desenvolvida com a linguagem Dart e o framework Flutter. O objetivo final é fortalecer a OBAMA como uma ferramenta indispensável para o planejamento docente.

Metodologia

Para atingir o objetivo de transformar a Busca Avançada em um serviço mais eficiente, a metodologia adotada foi centrada na experiência do usuário (UX) e na reorganização visual, utilizando Dart e Flutter como base tecnológica:

A) Análise da Usabilidade da Interface Anterior: O primeiro passo foi avaliar a disposição e acessibilidade dos filtros na versão anterior, identificando pontos de confusão ou dificuldade de acesso que pudessem atrapalhar o usuário.

B) Design de Interface com Foco no Fluxo do Usuário: Com base na análise, optou-se por implementar um sistema de busca em modal. Esta escolha de design é fundamental porque:

Contextualiza a Busca: O modal aparece sobre a listagem de OAs, permitindo que o professor refine sua busca sem perder o contexto da página em que estava, um princípio de usabilidade defendido por Cybis et al. (2007).

Agrupar Logicamente: Todos os filtros pedagógicos foram reunidos em um único local, organizados em seções claras, facilitando o scan visual e a seleção.

Torna a Busca uma Ação Rápida: Abrir, preencher e aplicar os filtros torna-se um processo fluido e rápido, encorajando a experimentação com diferentes combinações.

C) Implementação e Integração: A nova interface foi desenvolvida com Dart e Flutter como parte do processo de modernização da plataforma. A implementação focou em criar um componente frontend que capturasse as escolhas do usuário (os filtros selecionados) e as entregasse de forma clara para o sistema realizar a consulta ao banco de dados

Resultados e Discussões

O principal resultado deste trabalho foi a consolidação da Busca Avançada como um serviço robusto e user-friendly. A nova interface, desenvolvida em Flutter, centraliza e organiza a funcionalidade, representando uma evolução significativa em relação ao modelo anterior. Para contextualizar essa melhoria, a Figura 2 apresenta uma comparação side-by-side entre a implementação antiga da busca avançada e a nova proposta desenvolvida neste trabalho, destacando visualmente os avanços em termos de organização lógica, clareza visual e usabilidade.

Figura 1. Interface anterior | Figura 2. Nova interface

Os resultados alcançados são:

Otimização do Tempo do Professor: A nova organização permite que um professor localize um OA adequado em poucos cliques, filtrando diretamente por ano, tema e habilidade, o que é crucial para o seu cotidiano atribulado.

Precisão nos Resultados: A interface redesenhada facilita o uso de combinações de filtros, resultando em uma listagem de OAs muito mais aderente à necessidade específica do que uma busca por palavras-chave.

Promoção da Exploração Pedagógica: Ao tornar os filtros mais visíveis e acessíveis, a nova interface incentiva o professor a explorar recursos para outros anos ou temas relacionados descobrindo novas possibilidades para suas aulas.

Organização como Disseminação: Uma ferramenta bem organizada e fácil de usar é uma ferramenta que será adotada e recomendada. Portanto, este desenvolvimento contribui diretamente para o objetivo de disseminar a plataforma, pois reduz a frustração e aumenta a satisfação e a eficiência do usuário final.

Conclusão

Este trabalho demonstrou que o desenvolvimento de uma plataforma educacional vai muito além de compilar recursos, trata-se de criar serviços inteligentes que facilitem o acesso a eles. O redesenho da Busca Avançada para Plataforma OBAMA focou em organizar uma funcionalidade já existente sob uma nova perspectiva centrada no usuário, transformando-a em uma ferramenta estratégica para o professor.

A reorganização dos filtros em uma interface modal moderna desenvolvida com Dart e Flutter não apenas cumpriu o objetivo de organizar as funcionalidades da plataforma, mas também elevou o padrão de usabilidade, tornando a busca por recursos um processo intuitivo, rápido e alinhado com o raciocínio pedagógico. Investir na qualidade da busca é investir na própria efetividade da plataforma como instrumento de inovação na educação.

Referências

Batista, S., Bezerra, V., Oliveira, A. & Maia, D. (2017). OBAMA: Um repositório de objetos de aprendizagem para matemática. In Anais dos Workshops do Congresso Brasileiro de Informática na Educação, 300. DOI: <https://doi.org/10.5753/cbie.webie.2017.300>.

Bezerra, V., Batista, S., Oliveira, A., & Maia, D. (2019). Processo de identificação e correção de erros da plataforma OBAMA. In Anais do Workshop de Informática na Escola (pp. 1384–1388). SBC. <https://sol.sbc.org.br/index.php/wie/article/download/13322/13175/>.

Cybis, W.; Betiol, A. H.; Faust, R. (2007). Ergonomia e Usabilidade: Conhecimentos, Métodos e Aplicações. São Paulo: Editora Novatec.

Oliveira, A., Batista, S., Nascimento, Í., Azevedo, D., Lima, R., Oliveira, N. & Maia, D. (2018). Processo de desenvolvimento de uma ferramenta destinada à elaboração de planos de aula de forma colaborativa. In Anais do III Congresso sobre Tecnologias na Educação (CTRL+E 2018). UFC. https://ceur-ws.org/Vol-2185/CtrlE_2018_paper_118.pdf.

Anexos

figura 1

figura 2

Modal

CÓDIGO: HS1760

AUTOR: EMANUEL KYWAL PINTO CABRAL FILHO

ORIENTADOR: DENNYS LEITE MAIA

TÍTULO: Potimap: uma ferramenta interativa para análise do Estado do Rio Grande do Norte

Resumo

Os mapas sempre desempenharam papel essencial na sociedade, auxiliando na organização do espaço, na formulação de políticas e na compreensão das dinâmicas sociais e territoriais. No Brasil e no Rio Grande do Norte, sua evolução acompanha a própria formação histórica e administrativa. Nesse contexto, este trabalho apresenta o desenvolvimento da Potimap, uma ferramenta voltada à análise territorial do estado do RN, concebida para auxiliar na interpretação das divisões municipais, microrregionais e mesorregionais. A ferramenta foi construída a partir de um arquivo SVG otimizado e manipulado com a biblioteca SVG.js, que possibilita interações simples e eficientes. Entre as funcionalidades implementadas destacam-se: destaque visual de áreas e possibilidade de integração com conjuntos externos de informações por meio de identificadores do IBGE. Os resultados evidenciam que a solução atende às necessidades do Observatório Rede STEAM Potiguar, oferecendo leveza, simplicidade e escalabilidade. Embora apresente limitações relacionadas à acessibilidade e à ausência de análises geoespaciais avançadas, a Potimap se consolida como uma base promissora para evolução futura, incluindo publicação como pacote npm e integração a painéis dinâmicos. Assim, a ferramenta contribui para fortalecer a relação entre análise de dados potiguares e sua distribuição pelo território.

Palavras-chave: Ferramenta, Mapa, Rio Grande do Norte, Potiguar, SVG, STEAM

TITLE: Potimap: an interactive tool to analyse the State of Rio Grande do Norte

Abstract

Maps have always played an essential role in society, supporting spatial organization, policy-making, and the understanding of social and territorial dynamics. In Brazil and in the state of Rio Grande do Norte, their evolution reflects the historical and administrative formation of the territory. In this context, this paper presents the development of Potimap, a tool designed for territorial analysis of Rio Grande do Norte, aimed at assisting in the interpretation of municipal, micro-regional, and meso-regional divisions. The tool was built from an optimized SVG file and manipulated with the SVG.js library, enabling simple and efficient interactions. Among its main features are visual highlighting of areas and the possibility of integrating external datasets through IBGE identifiers. The results show that the solution meets the needs of the Observatório Rede STEAM Potiguar, offering lightness, simplicity, and scalability. Although it presents limitations related to accessibility and the lack of advanced geospatial analysis, Potimap establishes itself as a promising foundation for future development,

including publication as an npm package and integration with dynamic dashboards. Thus, the tool contributes to strengthening the relationship between data analysis in Rio Grande do Norte and its territorial distribution.

Keywords: Tool, map, Rio Grande do Norte, Potiguar, SVG, STEAM

Introdução

Desde os primórdios da civilização, os mapas constituem instrumentos essenciais para a organização da vida em sociedade. Eles não apenas representam o espaço físico, mas também orientam deslocamentos, estruturam fronteiras e influenciam decisões políticas, econômicas e culturais (Thrower, 2008).

Nesse sentido, compreender a distribuição espacial de indicadores sociais, econômicos e educacionais é fundamental para a criação de estratégias que favoreçam o desenvolvimento regional. Na atualidade, a cartografia digital amplia ainda mais esse papel, ao permitir a integração de dados dinâmicos e análises complexas que favorecem a compreensão de fenômenos sociais e territoriais (Carneiro, Sommer, 2017).

No Brasil, a história da cartografia acompanha a própria formação do país. Desde o período colonial, mapas foram elaborados para delimitar terras e consolidar o processo de ocupação territorial (Higa, Anzai, 2014). Com o avanço das técnicas de representação e a modernização dos órgãos oficiais, a cartografia brasileira se tornou ferramenta estratégica para o planejamento urbano, a gestão de recursos naturais e a formulação de políticas públicas (Linhares, Cobo, 2025).

No Rio Grande do Norte, os mapas desempenharam papel central na identificação de limites municipais, na administração regional e na condução de políticas de desenvolvimento. As representações cartográficas do estado, inicialmente restritas a registros estáticos, evoluíram para formatos digitais capazes de apoiar diagnósticos em diversas áreas, desde a agricultura até a educação. Essa evolução amplia a relevância de conhecer e interpretar o território potiguar de forma precisa e acessível.

Compreender o território do RN é fundamental para reconhecer suas diversidades regionais, desafios e potencialidades. O estado é composto por municípios, microrregiões e mesorregiões com características sociais, econômicas e culturais distintas. Ter acesso a ferramentas que permitam visualizar e analisar essas divisões fortalece a capacidade de gestores, pesquisadores e cidadãos de interpretar dados e formular estratégias adequadas às realidades locais.

Nesse sentido, mapas digitais interativos surgem como instrumentos de apoio à compreensão das características geopolíticas do território potiguar. Diferentemente dos mapas tradicionais, que apenas registram os limites geográficos, as soluções interativas possibilitam associar informações contextuais e promover análises comparativas. Assim, os mapas deixam de ser apenas representações estáticas e passam a constituir ambientes ativos de exploração e aprendizado.

Foi a partir dessa percepção que nasceu a proposta de desenvolver uma ferramenta voltada ao estado do Rio Grande do Norte, baseada em tecnologia vetorial e acessível por meio de navegadores. A ideia é permitir a manipulação direta dos elementos territoriais, destacando

regiões, exibindo dados vinculados e oferecendo novas formas de interação com o espaço geográfico potiguar.

A escolha pelo uso do formato SVG (Scalable Vector Graphics, ou Vetor Gráfico Escalável em tradução livre) aliado à biblioteca SVG.js buscou conciliar simplicidade técnica e eficiência prática. Além de ser um padrão aberto e amplamente suportado pelos navegadores modernos, o SVG possibilita a representação precisa dos limites territoriais sem perda de qualidade e permite que cada unidade geográfica seja tratada como elemento independente.

Dentro desse contexto, a aplicação inicial da ferramenta está diretamente ligada ao Observatório Rede STEAM Potiguar. Essa aplicação tem como missão auxiliar no desenvolvimento de um ecossistema de práticas criativas fundamentadas na abordagem STEAM, uma perspectiva pedagógica que busca estimular a criatividade, a interdisciplinaridade e o pensamento crítico. Ao integrar a nova solução cartográfica, o Observatório passa a contar com uma forma de visualizar territorialmente seus dados, fortalecendo análises e contribuindo para a difusão de práticas educacionais inovadoras

Metodologia

Este trabalho enquadra-se como uma pesquisa de desenvolvimento, pois concentra-se na concepção e implementação de uma ferramenta de software voltada à análise de dados a partir de uma perspectiva da geografia potiguar, privilegiando o processo de construção do artefato.

O desenvolvimento da Potimap partiu da combinação entre dois elementos centrais: (i) um arquivo SVG representando o mapa do estado do Rio Grande do Norte, devidamente estruturado em camadas de municípios, microrregiões e mesorregiões; e (ii) a utilização da biblioteca SVG.js, escolhida por sua simplicidade, curva de aprendizado reduzida e suporte direto a manipulações vetoriais. Essa escolha buscou equilibrar baixo custo de implementação com alto potencial de expansão, evitando a complexidade de bibliotecas mais robustas, como D3.js ou Mapbox, que exigiriam configurações adicionais não compatíveis com o desenvolvimento de um MVP (Minimum Viable Product, ou Produto Mínimo Viável).

O processo metodológico foi dividido em etapas: (1) preparação e otimização do arquivo SVG, (2) definição das interações desejadas e (3) implementação das funcionalidades com apoio do SVG.js.

2.1 Preparação do arquivo SVG

O mapa base do RN foi obtido a partir da Wikipédia, tendo sido a única fonte onde foi possível encontrar o mapa em formato SVG e com as divisões municipais, microrregionais e mesorregionais, como pode ser visto na figura 1.

Figura 1 (Raphael Lorenzeto de Abreu, Wikipédia, 2006)

A partir da imagem original (figura 1) e do Figma, uma ferramenta de design que permite criação e edição de arquivos vetorizados, foram realizados os seguintes tratamentos:

Otimização do arquivo: simplificação e organização de elementos complexos e remoção de elementos desnecessários.

Identificação semântica: atribuição de identificadores (id) e metadados (data-ibge, data-nome) a cada unidade territorial, assegurando a vinculação posterior com conjuntos de dados.

A adaptação gerada a partir desses tratamentos pode ser vista logo abaixo, na figura 2.

Figura 2 (Adaptado de Raphael Lorenzeto de Abreu, Wikipédia, 2006)

2.2 Definição das interações

O mapa não deveria ser apenas um objeto estático, mas sim uma ferramenta interativa. Assim, as interações previstas foram:

Seleção do tipo de divisão territorial que será interagido: municípios, microrregiões e mesorregiões.

Seleção do tipo de divisão territorial que será visualizado, podendo visualizar cada um dos três tipos individualmente ou em conjunto.

Selecionar unidades territoriais por meio de cliques, com suporte a seleção múltipla.

Alterar dinamicamente estilos visuais (cores, opacidade, bordas) de acordo com os dados vinculados.

2.3 Implementação com biblioteca SVG.js

A escolha pelo SVG.js se mostrou determinante para a simplicidade do desenvolvimento. A biblioteca oferece uma API intuitiva para criar, manipular e animar elementos SVG, o que permitiu:

Carregamento e manipulação do SVG: o mapa foi importado diretamente e controlado pelo objeto principal do SVG.js, possibilitando a seleção dos elementos (`SVG('#id')`) e a aplicação de métodos como `.fill()`, `.stroke()` e `.opacity()`.

Eventos de interação: cada elemento (município, microrregião e mesorregião) recebeu listeners nativos da biblioteca, como `.mouseover()`, `.mouseout()` e `.click()`, assegurando respostas imediatas às ações do usuário.

Animações e destaques: o SVG.js oferece recursos de animação nativos, permitindo transições suaves de cores, zoom em áreas específicas e mudanças graduais de opacidade. Essas animações foram empregadas para tornar a experiência mais fluida e intuitiva.

Resultados e Discussões

A implementação da ferramenta interativa descrita neste trabalho resultou em um protótipo funcional que cumpre os objetivos estabelecidos na introdução e na metodologia. O uso da biblioteca SVG.js possibilitou o desenvolvimento de uma interface capaz de manipular diretamente os elementos vetoriais do mapa do Rio Grande do Norte (RN), garantindo respostas rápidas às interações e permitindo a vinculação de dados externos de forma simplificada. Nesta seção, são apresentados os resultados alcançados, as dificuldades enfrentadas, as vantagens observadas em relação a outras soluções e as discussões sobre as implicações dessa abordagem no contexto do Observatório Rede STEAM Potiguar e em cenários mais amplos de análise territorial.

3.1 Resultados obtidos na implementação

O produto final, ainda em sua versão inicial, possibilita a realização das seguintes operações:

Exibição do mapa do RN em três níveis de granularidade (municípios, microrregiões e mesorregiões).

Destaque visual de regiões por meio de interações do usuário, com aplicação dinâmica de cores e efeitos de opacidade.

Integração de dados externos (JSON ou CSV) a partir de chaves IBGE, vinculando informações contextuais aos elementos gráficos do mapa.

3.2 Limitações encontradas

Apesar dos avanços, algumas limitações também foram identificadas:

Dependência do SVG.js: embora seja leve e simples, a biblioteca pode não oferecer suporte nativo para funcionalidades mais avançadas, como análise espacial ou sobreposição de camadas complexas.

Acessibilidade: a versão atual ainda carece de responsividade e mecanismos robustos de acessibilidade digital (ex.: suporte pleno a leitores de tela e navegação por teclado).

Funcionalidades de análise limitadas: a primeira versão prioriza a visualização e seleção, mas ainda não implementa cálculos estatísticos ou filtros dinâmicos diretamente no mapa.

3.3 Discussão comparativa

Ao comparar a ferramenta desenvolvida com alternativas já existentes, alguns pontos podem ser destacados:

D3.js: oferece maior poder de personalização e suporte a análises complexas, mas exige curva de aprendizado mais acentuada e código mais verboso. Para a primeira versão da ferramenta, isso representaria um obstáculo.

Leaflet e Mapbox: são mais adequadas para mapas interativos em escala global, baseados em tiles, mas pouco eficientes para manipulação de arquivos SVG locais e recortes regionais específicos.

SVG.js: se mostrou mais alinhada com os objetivos de simplicidade, rapidez e foco em interações básicas sobre elementos vetoriais já definidos.

Essa comparação evidencia que a escolha metodológica foi adequada ao contexto e ao estágio atual do Observatório, embora não descarte a possibilidade de, no futuro, integrar a ferramenta a bibliotecas mais robustas para atender demandas analíticas mais sofisticadas.

3.4 Perspectivas de evolução

Com base nos resultados alcançados e nas discussões levantadas, destacam-se alguns caminhos para evolução:

Publicação da ferramenta como pacote npm, permitindo que outros projetos façam uso do recurso.

Integração com painéis dinâmicos de análise, possibilitando filtros temporais, séries históricas e cruzamento de múltiplos indicadores em tempo real.

Incorporação de técnicas de análise de dados mais avançadas através da biblioteca D3.js.

Conclusão

O desenvolvimento da ferramenta Potimap demonstrou a viabilidade de criar uma solução interativa, simples e escalável para análise territorial do estado do Rio Grande do Norte. A escolha pelo uso de arquivos SVG em conjunto com a biblioteca SVG.js mostrou-se acertada, pois permitiu implementar interações básicas — como destaque e vinculação de dados externos — de forma intuitiva e com baixo custo de implementação. Esses resultados comprovam que a abordagem adotada atende às necessidades imediatas do Observatório Rede STEAM Potiguar, ao mesmo tempo em que estabelece uma base sólida para evoluções futuras.

Entre as principais contribuições da ferramenta, destacam-se: (i) a possibilidade de representar o mapa do RN em diferentes níveis de granularidade (municípios, microrregiões e mesorregiões) e (ii) a integração facilitada a partir da identificação dos municípios, microrregiões e mesorregiões a partir dos identificadores utilizados pelo IBGE.

Por outro lado, as limitações identificadas também apontam caminhos importantes para evolução. Aspectos como acessibilidade digital, escalabilidade para mapas maiores e integração com técnicas avançadas de análise geoespacial permanecem como desafios a serem enfrentados.

De forma geral, a ferramenta abre espaço para futuras expansões e colaborações, contribuindo não apenas para a análise territorial do RN, mas também para a discussão mais ampla sobre o uso de tecnologias digitais na educação.

Referências

CARNEIRO, E. L.; SOMMER, J. A. P. Cartografia temática digital. Canoas: Universidade Luterana do Brasil – ULBRA, 2017.

HIGA, T. C. C. de S.; ANZAI, L. C. Considerações sobre a cartografia colonial do Brasil: Categorias tipológicas. Revista do Instituto Histórico e Geográfico de Mato Grosso, Cuiabá (MT), v. 1, n. 74, p. 13–30, 2014. Disponível em: <https://revistaihgmt.com.br/index.php/revistaihgmt/article/view/331>. Acesso em: 31 ago. 2025.

LINHARES, L.; COBO, B. Rede de cidades e desenvolvimento regional no Brasil: uma leitura econômico-espacial em duas escalas. Revista Brasileira de Geografia, v. 69, p. 3-17, 2025. DOI:https://doi.org/10.21579/issn.2526-0375_2024_n2_3-17.

THROWER, N. J. W. Maps and Civilization: Cartography in Culture and Society. 3. ed. Chicago: University of Chicago Press, 2008.

Anexos

Figura 1 (Raphael Lorenzeto de Abreu, Wikipédia, 2006)

Figura 2 (Adaptado de Raphael Lorenzeto de Abreu, Wikipédia, 2006)

CÓDIGO: SB0309

AUTOR: ANA BEATRIZ CAMILO DA COSTA

ORIENTADOR: JORGE ESTEFANO DE SANTANA SOUZA

TÍTULO: Desenvolvimento de Soluções com Microserviços e Containers para Análises de Genomas Mitocondriais

Resumo

Este projeto teve como objetivo principal criar aplicações de suporte que facilitam análises de bioinformática, através de uma solução de software modular baseada em arquitetura de microservices e containers. Isso proporcionará escalabilidade, resiliência, autogerenciamento e manutenção aprimorada. Foram usadas metodologias ágeis para organizar, delegar e executar as tarefas de desenvolvimento. O foco atual foi desenvolvimento de soluções de software integradas à infraestrutura computacional existente, através de interfaces de usuário que incentivem a criação de ambientes propícios ao desenvolvimento e pesquisa.

Palavras-chave: Bioinformática; Alinhamento Genético; Análise de Genomas; Pipelines;

TITLE: Development of Solutions with Microservices and Containers for Mitochondrial Genome Analysis

Abstract

The main objective of this project was to develop support applications that facilitate bioinformatics analyses through a modular software solution based on a microservices and container architecture. This approach provides scalability, resilience, self-management, and improved maintainability. Agile methodologies were employed to organize, delegate, and execute development tasks. The current focus has been on developing software solutions integrated with the existing computational infrastructure, through user interfaces that foster the creation of environments conducive to development and research.

Keywords: Bioinformatics; Genetic Alignment; Genome Analysis; Pipelines;

Introdução

O Ambiente Multidisciplinar de Bioinformática (BioME) é essencial na formação de profissionais em bioinformática, impactando tanto a academia quanto os setores produtivos. Com equipes de pesquisa multidisciplinares, o BioME tem promovido avanços significativos em várias áreas da bioinformática, como biologia do câncer, modelagem de sistemas, genômica, proteômica, evolução molecular e bioinformática estrutural. No âmbito do Instituto Metrópole Digital (IMD) da UFRN, o BioME busca impulsionar o desenvolvimento da bioinformática.

A análise de grandes volumes de dados biológicos, como os terabases de sequências brutas utilizados na montagem de mitogenomas, representa um dos principais desafios atuais da bioinformática. A complexidade inerente a esses processos demanda soluções computacionais capazes de lidar não apenas com o alto custo de processamento, mas também com a necessidade de garantir reprodutibilidade, escalabilidade e eficiência. Nessa perspectiva, a qualidade, a clareza e a organização do código tornam-se elementos fundamentais para o sucesso das análises.

Com esse entendimento, o presente projeto concentrou-se na revisão crítica e na refatoração de partes do código empregado em estudos anteriores. O objetivo principal foi compreender em profundidade a lógica implementada e, a partir disso, reestruturar seções que apresentavam fragilidades relacionadas à legibilidade, modularidade e organização. Essa abordagem buscou transformar trechos de código excessivamente complexos ou fragmentados em componentes mais claros, coesos e facilmente mantidos.

O trabalho proposto vai além da simples correção técnica: ele contribui para a consolidação de uma base computacional mais robusta e confiável, capaz de sustentar pipelines de bioinformática de maior escala e complexidade. Ao adotar princípios de desenvolvimento de software modular, alinhados às melhores práticas de engenharia de software, reforça-se não apenas a eficiência imediata dos processos, mas também a possibilidade de adaptação e expansão futura. Dessa forma, este projeto representa um passo importante para viabilizar análises mais consistentes de genomas mitocondriais, ampliando as condições para estudos de diversidade genética em espécies sub-representadas, especialmente em ecossistemas de alta relevância biológica, como a Amazônia.

Desta forma o objetivo deste plano de trabalho é continuar o desenvolvimento de uma solução de software modular baseada em arquitetura de microservices e containers, oferecendo maior escalabilidade, resiliência, autogerenciamento e manutenção. O projeto, direcionado para a análise de mitocôndrias de espécies endêmicas da região amazônica, pretende escalar a extração e o processamento de dados biológicos, tornando essas tarefas mais eficientes para a comunidade do BioME/IMD/UFRN assim como para a comunidade científica.

Metodologia

O desenvolvimento e a validação do software foram conduzidos a partir da adaptação e análise crítica do pipeline wf-alignment, disponibilizado em código aberto pela Oxford Nanopore Technologies Plc no GitHub (<https://github.com/epi2me-labs/wf-alignment>). Esse pipeline é construído em Nextflow e integra diferentes ferramentas consolidadas na área de bioinformática, entre elas:

minimap2 – utilizado para o alinhamento eficiente de leituras curtas e longas contra sequências de referência;

samtools – aplicado ao processamento, manipulação e indexação de arquivos no formato BAM;

mosdepth – empregado para o cálculo da profundidade de cobertura ao longo das sequências de referência;

seqtk – usado para manipulação básica de sequências em formatos FASTA/FASTQ.

As leituras de entrada foram simuladas e processadas em formatos FASTQ e BAM não alinhados. Como referência, foram selecionados genomas mitocondriais de espécies-modelo, obtidos diretamente do banco de dados NCBI GenBank (<https://www.ncbi.nlm.nih.gov/genbank/>).

O fluxo de processamento consistiu em:

Preparação e organização das sequências de referência em diretórios específicos para entrada no pipeline;

Alinhamento das leituras contra os genomas de referência utilizando minimap2;

Conversão e indexação dos resultados em formato BAM com o auxílio do samtools;

Cálculo da profundidade de cobertura por meio do mosdepth;

Geração de estatísticas de alinhamento e relatórios sumarizados pelo wf-alignment.

Refatoração de código

Durante a aplicação prática do pipeline, foram identificadas limitações relacionadas ao processamento de arquivos de grande porte e à escalabilidade do fluxo em cenários de alta demanda computacional. Para mitigar esses problemas, realizou-se uma refatoração de trechos específicos do código. Essa etapa envolveu:

Revisão sistemática do fluxo em Nextflow, visando reorganizar processos redundantes e otimizar a estrutura de dependências;

Modularização de etapas críticas, permitindo maior clareza e reuso do código em diferentes contextos de análise;

Simplificação de scripts auxiliares, reduzindo a complexidade lógica e aumentando a legibilidade;

Implementação de logs e mensagens de erro mais informativas, para facilitar a depuração e a manutenção futura.

Metodologias ágeis

O processo de desenvolvimento seguiu princípios de metodologias ágeis, em especial práticas inspiradas no Scrum e no Kanban. As tarefas foram organizadas em ciclos curtos de desenvolvimento (sprints), priorizadas em função dos gargalos identificados e acompanhadas em quadros de gerenciamento de fluxo. Essa abordagem permitiu:

Maior rapidez na identificação e correção de problemas;

Incrementos graduais na eficiência e manutenibilidade do pipeline;

Documentação contínua das mudanças realizadas, garantindo rastreabilidade;

Flexibilidade para ajustar o desenvolvimento conforme surgiam novos desafios técnicos.

Essa combinação de análise bioinformática, refatoração de código e métodos ágeis de desenvolvimento resultou em um pipeline mais claro, escalável e robusto, adequado para aplicações futuras em análises genômicas de maior complexidade.

Resultados e Discussões

Até o momento, o software em desenvolvimento encontra-se em fase de testes e já é capaz de realizar a leitura e o processamento de arquivos de entrada nos formatos FASTQ e BAM, executar o alinhamento das sequências genéticas contra os genomas de referência selecionados e gerar estatísticas descritivas do processo de alinhamento. Essas estatísticas incluem métricas como número total de leituras alinhadas, percentual de mapeamento, distribuição de erros e profundidade de cobertura ao longo das sequências de referência.

Os resultados preliminares indicam que o sistema é funcional e eficiente em cenários com conjuntos de dados de pequeno a médio porte, apresentando boa consistência nas métricas geradas e reprodutibilidade entre execuções. Entretanto, foram observadas limitações no processamento de arquivos de maior volume, particularmente no tempo de execução e no uso de memória. Esses gargalos dificultam a escalabilidade do pipeline em análises que envolvem bancos de referência mais extensos ou conjuntos massivos de leituras de entrada.

A análise exploratória sugere que tais restrições estão associadas tanto ao desenho original do pipeline quanto à forma como as dependências entre processos são gerenciadas no Nextflow. Alguns módulos apresentaram acúmulo de processos intermediários, o que pode impactar a eficiência em cenários de alto paralelismo. Além disso, a manipulação de arquivos temporários de grande porte mostrou-se um fator crítico para a estabilidade do sistema.

Como resposta a essas limitações, foram conduzidas etapas iniciais de refatoração de código, que resultaram em melhorias na clareza e na modularidade do pipeline. A reorganização de etapas redundantes e a introdução de mecanismos de monitoramento de recursos permitiram identificar pontos críticos do fluxo de execução, preparando o terreno para otimizações futuras. Ainda assim, a investigação das causas subjacentes aos gargalos de desempenho permanece em andamento, com foco na mitigação de falhas potenciais e no aumento da robustez e escalabilidade do sistema.

De forma geral, os resultados obtidos até aqui validam a viabilidade da ferramenta para aplicações em bioinformática, mas também reforçam a necessidade de ajustes e aprimoramentos contínuos. O projeto encontra-se em uma etapa intermediária, em que a base funcional já está estabelecida, mas a otimização de desempenho e a

adaptação para grandes volumes de dados ainda são objetivos centrais a serem alcançados.

Conclusão

Durante o desenvolvimento do projeto de iniciação científica, a bolsista adquiriu conhecimentos aprofundados sobre estudos genômicos e ampliou sua compreensão de processos moleculares, a partir de referências bibliográficas especializadas. Paralelamente, obteve experiência prática no desenvolvimento de pipelines em Nextflow, assim como na criação e gestão de imagens e containers em Docker, habilidades essenciais para o manuseio e a escalabilidade de workflows de bioinformática.

Em termos de resultados do projeto, a pesquisa concentrou-se no desenvolvimento e refatoração de um software capaz de realizar alinhamentos genéticos e análise de dados de genomas mitocondriais. Embora os resultados preliminares demonstrem eficiência no processamento de dados e geração de estatísticas, as conclusões finais sobre a precisão e a robustez do sistema permanecem em avaliação. Esse processo contínuo de validação é fundamental para garantir a veracidade estatística das análises genômicas e apoiar estudos futuros.

O trabalho realizado evidencia a necessidade de sistemas computacionais bem estruturados e escaláveis para aumentar a eficiência e a confiabilidade das análises de alinhamento genético. Além de contribuir para o avanço técnico do projeto, a experiência proporcionou uma base sólida para pesquisas futuras, reforçando a importância de práticas de desenvolvimento de software modular e pipelines reproduzíveis no estudo de genomas.

Referências

- 1- Fowler, M. (2014). Microservices. martinfowler.com.
- 2- Beck, K., et al. (2001). Manifesto for Agile Software Development. Agile Alliance.
- 3- Newman, S. (2015). Building Microservices. O'Reilly Media.
- 4- Andrews, R. M., et al. (1999). Reanalysis and revision of the Cambridge reference sequence for human mitochondrial DNA. Nature Genetics.
- 5- Django Software Foundation. Django: The Web framework for perfectionists with deadlines.
- 6- Facebook Inc. React - A JavaScript library for building user interfaces.
- 7- Evans, E. (2003). Domain-Driven Design: Tackling Complexity in the Heart of Software. Addison-Wesley Professional.
- 8- Humble, J., & Farley, D. (2010). Continuous Delivery: Reliable Software Releases through Build, Test, and Deployment Automation. Addison-Wesley Professional.

9- Seringhaus, M. R., & Gerstein, M. B. (2008). Genomics Confounds Gene Classification. American Scientist.

10-Altman RB, Dugan JM. Defining bioinformatics and structural bioinformatics. Methods Biochem Anal. 2003;44:3-14. PMID: 12647379.

CÓDIGO: SB0360

AUTOR: JOAREMIO MARINHO REVOREDO NETO

ORIENTADOR: JORGE ESTEFANO DE SANTANA SOUZA

TÍTULO: Desenvolvimento de Ferramentas para Análise de Sintenia e Colinearidade de Regiões Genômicas

Resumo

Este plano de trabalho teve como objetivo principal o estudo de ferramentas de montagem e processos relacionados a obtenção e anotação de Mitogenomas assim como a sua composição e ordenação genica. Através de uma abordagem de revisão de código e aplicação de princípios de refatoração, buscamos melhorar a clareza, a manutenibilidade e a eficiência de seções de pipelines de bioinformática que apresentavam complexidade ou organização subótima. Este trabalho visa aprimorar a fundação computacional para futuras análises de genomas mitocondriais, facilitando a replicação e a extensão da pesquisa em diversidade genética de espécies sob representadas da Amazônia.

Palavras-chave: Bioinformática; Genomas Mitocondriais; Refatoração de Código;

TITLE: Development of Tools for the Analysis of Synteny and Collinearity of Genomic Regions

Abstract

This work plan had as its main objective the study of assembly tools and processes related to the acquisition and annotation of mitogenomes, as well as their gene composition and organization. Through a code review approach and the application of refactoring principles, we sought to improve the clarity, maintainability, and efficiency of sections of bioinformatics pipelines that exhibited complexity or suboptimal organization. This work aims to enhance the computational foundation for future analyses of mitochondrial genomes, facilitating the replication and extension of research on the genetic diversity of underrepresented species from the Amazon.

Keywords: Bioinformatics; Mitochondrial Genomes; Code Refactoring

Introdução

A análise de genomas mitocondriais (mitogenomas) representa uma ferramenta fundamental para desvendar a biodiversidade, a genética populacional e a evolução das espécies. Suas aplicações são cruciais em áreas como a conservação, a gestão de pescas e o monitoramento ambiental. No entanto, espécies do Sul Global, e em particular as da vasta e biodiversa Amazônia, permanecem sub-representadas em bancos de dados genômicos públicos. Esta lacuna limita significativamente o conhecimento sobre a biodiversidade local e obstrui a implementação de estratégias de conservação eficazes.

O estudo "Diversidade Mitogenômica em Peixes Amazônicos: Montagem, Anotação e Evolução a Partir de 100 Novos Mitogenomas" abordou essa questão ao montar e caracterizar 100 novos mitogenomas de 34 espécies de peixes amazônicos, utilizando dados de sequenciamento de genoma completo (WGS) disponíveis publicamente. Este esforço gerou recursos genômicos inéditos para 22 espécies até então ausentes no GenBank, contribuindo substancialmente para preencher essas lacunas críticas e promover estratégias de conservação e gestão sustentável.

A complexidade inerente ao processamento de grandes volumes de dados biológicos, como os terabases de sequências brutas utilizados na montagem dos mitogenomas, demanda soluções de software robustas e eficientes. Para garantir a eficiência e a escalabilidade de tais análises, a qualidade e a organização do código subjacente são primordiais. Nesse contexto, este projeto concentrou-se na análise e refatoração de partes do código utilizado no estudo supracitado. Nosso objetivo foi entender profundamente a lógica implementada e aprimorar as seções do código que se encontravam "bagunçadas", visando aumentar a clareza, a manutenibilidade e a resiliência dos pipelines, em alinhamento com os princípios de desenvolvimento de software modular para bioinformática.

Para alcançar os objetivos de escalabilidade e eficiência na análise de genomas mitocondriais de espécies amazônicas, o desenvolvimento da solução de software modular integra-se a metodologias e ferramentas robustas. Uma dessas ferramentas centrais é o Nextflow, uma plataforma amplamente reconhecida por facilitar a criação de fluxos de trabalho computacionais reprodutíveis e escaláveis para análises genômicas. No contexto deste projeto, o software em desenvolvimento "tem influência das pipelines em Nextflow", sendo a base para funcionalidades críticas como a combinação de arquivos de referência, o alinhamento de leituras de entrada (provenientes de arquivos FASTQ ou BAM não alinhados) e a criação de estatísticas de alinhamento, montagem de genomas, anotação e análises de desintenia e colinearidade. Essa abordagem com Nextflow é fundamental para o processamento eficiente de dados e para a realização de inferências precisas sobre as procedências genéticas do material biológico em estudo. Além disso, a montagem inicial dos genomas mitocondriais de peixes amazônicos, conforme descrito em um estudo relacionado, foi especificamente realizada utilizando uma pipeline em desenvolvimento em Nextflow. Este uso do Nextflow sublinha a relevância da ferramenta para a escalabilidade da extração e processamento de dados biológicos, um dos pilares deste projeto.

Metodologia

Este trabalho baseou-se em uma metodologia de revisão e refatoração de código, aplicando princípios de engenharia de software à pipeline bioinformática utilizada na pesquisa de diversidade mitogenômica em peixes amazônicos. O processo envolveu as seguintes etapas:

- **Compreensão do Pipeline Existente:** Inicialmente, foi realizada uma análise detalhada do fluxo de trabalho e dos scripts utilizados para a montagem, anotação e análise dos mitogenomas e sua estruturação (sintenia e colinearidade). Isso incluiu a revisão das etapas de aquisição de dados, análise de qualidade, montagem de mtDNA (utilizando ferramentas como NOVOPlasty e Unicycler) e anotação (MitoAnnotator, MITOS2).
- **Identificação de Seções para Refatoração:** Durante a fase de compreensão, foram identificadas seções do código que apresentavam características de baixa clareza, duplicação de lógica ou organização subótima. Estas seções foram priorizadas para intervenção, com o objetivo de melhorar sua legibilidade e manutenção sem alterar seu comportamento externo.
- **Aplicação de Técnicas de Refatoração:** Para cada seção identificada, foram aplicadas técnicas de refatoração, tais como:
 - **Renomeação de variáveis e funções:** Para melhorar a semântica e a clareza do código.
 - **Extração de métodos/funções:** Para modularizar lógicas complexas e reduzir a duplicação.
 - **Simplificação de condicionais e loops:** Para tornar o fluxo de execução mais intuitivo.
 - **Remoção de comentários:** Remoção de linhas de código e comentários desnecessários para o projeto e estavam colocados de maneira desorganizada. Porém continuam salvos pelo gerenciamento de versões do código.

Resultados e Discussões

Até o momento atual, a fase de análise e refatoração do código-base da pesquisa "Diversidade Mitogenômica em Peixes Amazônicos" resultou em melhorias significativas na estrutura interna e na legibilidade dos scripts. Concretamente, observamos os seguintes resultados iniciais:

- **Aumento da Clareza e Legibilidade:** Seções do código que antes eram difíceis de interpretar devido à ausência de modularidade ou nomes pouco descritivos agora possuem uma estrutura mais lógica e nomes de variáveis/funções mais intuitivos. Isso facilita o entendimento do fluxo de processamento dos dados genômicos, desde a leitura de arquivos SRA até as análises de características mitogenômicas.

- Identificação de Oportunidades de Otimização Futuras: A compreensão aprofundada do código revelou pontos onde otimizações de desempenho podem ser aplicadas no futuro. Embora a refatoração inicial tenha focado na clareza, a análise permitiu mapear gargalos ou algoritmos que poderiam ser reescritos para maior eficiência computacional, especialmente considerando a vasta quantidade de dados processados.

Conclusão

Este projeto de refatoração e análise de código demonstra o valor intrínseco de investir na qualidade do software que suporta a pesquisa bioinformática de ponta.

A capacidade de revisar, entender e aprimorar o código existente é crucial para o avanço da bioinformática, permitindo que a ciência se construa sobre fundações computacionais mais robustas e verificáveis. Ao tornar o pipeline mais acessível e fácil de manter, este trabalho contribui indiretamente para a promoção de estratégias de conservação mais eficazes e o manejo sustentável da rica biodiversidade aquática da Amazônia, ao facilitar a pesquisa contínua e a aplicação de novas ferramentas para identificação e monitoramento de espécies.

Futuros esforços incluirão a implementação de otimizações de desempenho identificadas durante esta fase de refatoração, possivelmente explorando arquiteturas mais modulares como microserviços e containers para lidar com volumes de dados ainda maiores e garantir maior escalabilidade e resiliência. Em última análise, este trabalho reforça a necessidade de uma abordagem holística na pesquisa bioinformática, onde o rigor científico dos dados é complementado pela excelência na engenharia de software.

Referências

NEXTFLOW. NextflowDocumentation. Disponível em:
<https://www.nextflow.io/docs/latest/>.

SILVA, Larissa Martins Brito e; AZEVEDO, Gleison Medeiros de; FILHO, Júlio César da Silva; MEDEIROS, Iruziky Araujo de; GOMES, Daniel Henrique Ferreira; GUERREIRO, Sávio Lucas de Matos; DALMOLIN, Rodrigo Juliani Siqueira; SOUZA, Jorge Estefano de Santana. UnveilingMitogenomicDiversity in AmazonianFish: Assembly, Annotation, andEvolutionaryfrom 100 New Mitogenomes. [s.l.]: [s.n.], 2023.

BRASIL ESCOLA. O que é Genoma. Disponível em:
<https://brasilecola.uol.com.br/o-que-e/biologia/o-que-e-genoma.htm>.

BRASIL ESCOLA. Mitocôndrias. Disponível em:
<https://brasilecola.uol.com.br/biologia/mitocondrias.htm>.

ACCURATE. Clean Code. Blog Accurate, 2023. Disponível em:
<https://blog.accurate.com.br/clean-code/>.

MARTIN, Robert C. Clean Code: A Handbook of Agile Software Craftsmanship. 1. ed. Upper Saddle River: Prentice Hall, 2008.

CÓDIGO: SB0878

AUTOR: DANIEL NOBRE DE QUEIROZ PEREIRA

COAUTOR: MARIA LUIZA DINIZ DE SOUSA LOPES

COAUTOR: LARA EMILY OLIVEIRA SOUSA

COAUTOR: GLEYDSON TEOTONIO DO NASCIMENTO

ORIENTADOR: FREDERICO ARAUJO DA SILVA LOPES

TÍTULO: Desenvolvimento e aprimoramento de algoritmo de classificação de gravidade de queilite actínica através de aprendizado de máquina

Resumo

A queilite actínica (QA) é uma condição potencialmente maligna que afeta os lábios, causada pela exposição crônica à radiação ultravioleta (UV). A diversidade clínica das manifestações e a ausência de correlação consistente com o grau histopatológico dificultam o manejo clínico. Este estudo reporta os avanços obtidos no último ano na construção de um algoritmo de aprendizado de máquina para classificar as QAs em baixo e alto risco, utilizando dados tabulares clínicos e demográficos de pacientes. A pesquisa é fruto da colaboração entre os Departamentos de Odontologia e Tecnologia da Informação da UFRN. Nesta fase, consolidamos um pipeline robusto para dados tabulares, atingindo métricas superiores às versões iniciais (Acurácia=66,7%, $F1=0,63$, $AUC=0,664$, $AUPR=0,611$). Também estruturamos a pipeline para imagens clínicas com segmentação do lábio e lesões, atualmente em fase de aprimoramento. Apesar da ausência de integração completa entre os dois eixos, os resultados parciais indicam potencial de uso clínico da abordagem proposta.

Palavras-chave: Inteligência Artificial; Aprendizado de Máquina; Queilite Actínica; CI

TITLE: Development and improvement of an algorithm for classifying the severity of actinic cheilitis through machine learning

Abstract

Actinic cheilitis (AC) is a potentially malignant lip disorder caused by chronic ultraviolet (UV) exposure. Clinical variability and lack of consistent correlation with histopathological grading hinder clinical management. This paper reports the advances achieved over the last year in developing a machine learning algorithm to classify AC into low- and high-risk categories, based on patient demographic and clinical tabular data. The project results from collaboration between the Dentistry and Information Technology Departments at UFRN. In this phase, we consolidated a robust pipeline for tabular data, achieving improved metrics over the initial versions (Accuracy=66.7%, $F1=0.63$, $AUC=0.664$, $AUPR=0.611$). We also structured the image pipeline (lip segmentation and lesion detection), which remains under refinement. Although full integration of both axes has not yet been completed, partial results highlight the feasibility of the proposed clinical support approach.

Keywords: Artificial Intelligence; Machine Learning; Cheilitis.

Introdução

1. INTRODUÇÃO

A queilite actínica (QA) representa uma condição precursora de carcinoma de células escamosas labial, associada à exposição solar crônica. O desafio clínico decorre da grande variabilidade nas apresentações clínicas e da dificuldade em prever, a partir de sinais visíveis, o grau histopatológico da lesão (MEDEIROS et al., 2019; POITEVIN et al., 2017). Nesse contexto, a Inteligência Artificial (IA), e em particular o Aprendizado de Máquina (ML), tem se destacado como ferramenta promissora para análise de dados clínicos complexos e suporte ao diagnóstico (ESTEVA et al., 2017; AL KUWAITI et al., 2023).

No primeiro estudo, reportamos uma pipeline inicial baseada em Support Vector Machines (SVM), que obteve acurácia em torno de 50%. Embora limitada, essa etapa demonstrou a viabilidade do uso de ML na classificação de QA. No presente trabalho, relatamos os avanços do último ano, em especial na análise de dados tabulares, onde a introdução de técnicas de seleção de variáveis, ajuste de thresholds e busca de hiperparâmetros resultou em melhora significativa de desempenho (STOLTZFUS, 2011). Também descrevemos o desenho da pipeline de imagens clínicas, ainda em desenvolvimento, que futuramente será integrada à análise tabular.

Metodologia

2. METODOLOGIA

Este estudo foi estruturado com o objetivo de desenvolver e validar um modelo de aprendizado de máquina para a classificação de lesões de queilite actínica (QA) em níveis de risco (baixo e alto), utilizando tanto dados clínicos e demográficos tabulares quanto imagens clínicas da região labial. A metodologia contemplou a preparação do dataset, seleção de variáveis, definição de algoritmos candidatos, avaliação com diferentes métricas e o início da construção de uma pipeline multimodal integrando informações tabulares e imagéticas.

2.1. Revisão da literatura e justificativa das escolhas

Diversas técnicas de aprendizado de máquina vêm sendo exploradas em contextos médicos para classificação de dados clínicos e histopatológicos. Entre elas:

- Máquinas de Vetores de Suporte (SVM): amplamente utilizadas em classificações biomédicas, escolhidas como baseline no estudo inicial, mas limitadas na predição da classe minoritária (ESTEVA et al., 2017).
- Regressão Logística: modelo estatístico validado na área biomédica, valorizado pela interpretabilidade clínica e pelo cálculo de odds ratios (STOLTZFUS, 2011).
- Métodos baseados em árvores (ExtraTrees, Random Forest): aplicados para avaliação da importância das variáveis e redução de ruído (BABY et al., 2021).

- Redes neurais convolucionais (CNNs): utilizadas neste estudo especificamente para segmentação labial e detecção de lesões, dada sua eficácia em padrões espaciais complexos (JUBAIR et al., 2022).

Essa revisão orientou a definição do pipeline, composto por dois eixos independentes (tabular e imagético), planejados para futura integração.

2.2. Coleta e preparação dos dados tabulares

Os dados tabulares foram obtidos a partir de registros clínicos e histopatológicos de pacientes diagnosticados com QA no Departamento de Odontologia da UFRN. O dataset foi inicialmente organizado em um arquivo CSV com variáveis demográficas, clínicas e de hábitos.

- Pré-processamento: exclusão de colunas redundantes (registros internos, idade duplicada) e substituição de valores inconsistentes (999 $\rightarrow$ 0).
- Tratamento de valores ausentes: remoção de registros sem diagnóstico histopatológico confirmado pelo examinador 3.
- Codificação: variáveis categóricas convertidas em strings e posteriormente codificadas com LabelEncoder.
- Mapeamento de classes: diagnóstico histopatológico em duas classes — Baixo Risco (sem displasia, displasia leve, outros achados benignos) e Alto Risco (displasia moderada, severa ou carcinoma epidermoide).

A Figura 1 apresenta o fluxograma do eixo tabular.

2.3. Seleção de variáveis

Aplicamos um ExtraTreesClassifier com 100 estimadores para calcular a importância relativa das variáveis. Em seguida, o SelectFromModel foi utilizado com limiar mediano, resultando em 11 preditores principais: cor da pele, tabagismo, etilismo, descamação, perda de demarcação, áreas pálidas, mancha branca, placa branca, endurecimento labial, sintomatologia dolorosa e prurido.

A Figura 2 apresenta o gráfico de importância das variáveis.

2.4. Modelagem

Foram avaliados dois algoritmos principais:

- Support Vector Classifier (SVC-RBF): ajustado com GridSearchCV ($C \in \{0.1, 1, 10, 50\}$, $\gamma \in \{\text{scale}, 0.01, 0.1, 1\}$).
- Regressão Logística: ajustada com busca em $C \in \{0.01, 0.1, 1, 10\}$, regularização L2, $\text{max_iter}=2000$.

A Regressão Logística foi priorizada pelo equilíbrio entre desempenho e interpretabilidade clínica.

A Figura 3 apresenta a comparação entre os modelos avaliados.

2.5. Avaliação e métricas

A avaliação foi conduzida com diferentes métricas:

- Matriz de confusão (Figura 4).
- Curva ROC (Figura 7).
- Curva Precision-Recall (Figura 8).
- Curva de calibração.
- Coeficientes e odds ratios: interpretação clínica das variáveis.
- Importância por permutação.
- Varredura de thresholds: explorando pontos de decisão alternativos.

2.6. Eixo de imagens

As imagens clínicas da região labial foram coletadas no Departamento de Odontologia da UFRN e anotadas no software LabelMe, com marcações da região labial (ROI) e das lesões.

A pipeline foi estruturada em duas etapas:

- Segmentação binária do lábio.
- Segmentação/detecção de lesões multiclases.

Na primeira versão, foi utilizado um UNet; na segunda, GroupNorm + ComboLoss para lábio e DeepLabV3-ResNet50 para as lesões.

A Figura 5 mostra o fluxograma completo (tabular + imagens), e a Figura 6 apresenta a frequência das lesões mais comuns.

Resultados e Discussões

3. RESULTADOS

No eixo tabular, o melhor desempenho foi alcançado pela Regressão Logística, que obteve acurácia de 66,7%, F1-score de 0,632, ROC-AUC de 0,664 e AUPR de 0,611. Quando o limiar de decisão foi ajustado para 0,20, o modelo passou a apresentar recall de 1.0 para a classe de alto risco, garantindo que nenhum caso grave fosse perdido, mesmo que à custa de uma diminuição na precisão (0,463). Esse ajuste reforça uma estratégia de priorização da sensibilidade, aspecto particularmente relevante em cenários clínicos preventivos, nos quais a falha em identificar um caso de alto risco pode trazer consequências graves. Por outro lado, o classificador SVC-RBF apresentou desempenho substancialmente inferior, com acurácia de 54,8% e incapacidade de prever a classe minoritária, refletida em um F1-score de 0.000, ROC-AUC de 0.364 e AUPR de 0.396. Esses resultados demonstram a inadequação do SVC-RBF para este problema específico.

A análise dos coeficientes da Regressão Logística indicou que variáveis como sintomatologia dolorosa, cor da pele, perda de demarcação e endurecimento labial foram positivamente

associadas ao alto risco, enquanto etilismo e prurido apareceram como fatores de associação negativa. Tais achados corroboram evidências já relatadas na literatura odontológica (BRITO et al., 2019; JUBAIR et al., 2022). As curvas ROC e Precision-Recall reforçaram que o comportamento do modelo foi moderado, mas superior ao baseline anterior, enquanto a curva de calibração revelou desvios intermediários, sugerindo a necessidade de bases mais extensas e balanceadas para garantir maior robustez probabilística.

No eixo de imagens, a primeira versão da pipeline adotou uma arquitetura UNet para segmentação binária do lábio, utilizando um dataset composto por 139 pares de imagem e arquivo JSON. Após 10 épocas de treinamento, observou-se convergência estável, com o train loss reduzindo de 0.49 para 0.21 e o validation loss de 0.55 para 0.18, resultando em predições satisfatórias da região labial. Paralelamente, foi montado um dataset com 101 imagens contendo máscaras válidas de lesões multiclases, após a padronização de rótulos inconsistentes. As lesões mais frequentes incluíram perda de demarcação (283 casos), mancha branca (272), área pálida (271), mancha vermelha (176), descamação (163) e placa branca (162). Esse conjunto serviu de base para o treinamento de um UNet multiclases, que apresentou diminuição progressiva do erro de treinamento ($2.76 \rightarrow 0.47$), embora o erro de validação tenha se estabilizado em 0.53, evidenciando limitações de generalização. Visualizações qualitativas mostraram sobreposições entre predições e máscaras de referência, ainda que com ativações limitadas em algumas classes.

Na segunda versão, o módulo de lábio foi aprimorado com o uso de GroupNorm, dropout2d e ComboLoss (BCE + Dice), treinado com AdamW e ReduceLROnPlateau. Houve melhora progressiva até aproximadamente a 15ª época, quando a validação estabilizou em torno de 0.39. Além disso, foram geradas métricas quantitativas por imagem, incluindo IoU, Dice, Precisão e Recall, organizadas em planilhas para avaliação mais detalhada. Para o módulo de lesões, foi implementado o DeepLabV3-ResNet50 ajustado para 18 classes, utilizando uma combinação de funções de perda e pesos proporcionais à frequência dos pixels. Apesar do setup robusto, os resultados não avançaram de forma satisfatória: o train loss reduziu de 1.14 para 0.55, mas a validação permaneceu instável entre 1.34 e 1.52, levando ao early stopping no nono ciclo, sem ganhos além do primeiro checkpoint. Esse comportamento reflete o sobreajuste e a dificuldade do modelo em lidar com classes de baixa frequência, cujos pesos chegaram a ordens de magnitude da ordem de 10^3 a 10^6 .

De maneira geral, os resultados sugerem que o módulo de lábio encontra-se funcional e estável, enquanto o módulo de lesões ainda apresenta limitações relacionadas ao desbalanceamento e à confiabilidade das predições. Assim, neste estágio, a integração entre os eixos tabular e imagético se encontra inviável no momento atual, sendo considerada uma meta de médio prazo.

4. DISCUSSÃO

Os resultados obtidos neste estudo reforçam a aplicabilidade do aprendizado de máquina no apoio ao diagnóstico clínico da queilite actínica. O uso da regressão logística, que apresentou desempenho superior ao baseline reportado no trabalho anterior, confirma a importância de etapas como a seleção de variáveis e o ajuste de thresholds. O fato de o modelo ter atingido recall de 1.0 para a classe de alto risco demonstra sua utilidade em cenários preventivos, nos quais o não reconhecimento de casos graves representa um risco clínico elevado. Ainda que essa estratégia implique redução na precisão, ela se mostra válida dentro do contexto de triagem clínica, priorizando a segurança do paciente.

No eixo de imagens, a consolidação do módulo de lábio representa um avanço importante, pois garante a extração consistente da região anatômica de interesse. Entretanto, as dificuldades encontradas no módulo de lesões refletem desafios clássicos da área, como o desbalanceamento severo do dataset e a baixa representatividade de determinadas classes. Esses fatores prejudicam tanto a estabilidade do modelo quanto sua capacidade de generalização. A literatura em IA aplicada à estomatologia mostra resultados semelhantes, nos quais ganhos tangíveis foram observados em tarefas supervisionadas, mas fortemente dependentes da ampliação e balanceamento das bases de imagens (ESTEVA et al., 2017; JUBAIR et al., 2022). À luz desses achados, as perspectivas futuras do presente estudo incluem o aumento do dataset de imagens, o uso de técnicas de data augmentation e a exploração de arquiteturas mais sofisticadas, com o objetivo de viabilizar a integração entre os eixos tabular e imagético em uma pipeline multimodal. Essa integração poderá resultar, no futuro, em uma ferramenta prática, não invasiva e de baixo custo, com potencial real de aplicação no manejo clínico da QA.

Conclusão

5. CONCLUSÕES

Os achados deste estudo evidenciam o potencial do aprendizado de máquina como ferramenta de apoio à decisão clínica no manejo da queilite actínica. No eixo tabular, a Regressão Logística consolidou-se como o modelo mais adequado, alcançando acurácia de 66,7% e F1-score de 0,632, além de recall de 1.0 para a classe de alto risco quando aplicado um threshold ajustado. Esse resultado reforça a relevância da seleção criteriosa de variáveis e do ajuste de parâmetros como estratégias fundamentais para garantir maior sensibilidade, característica essencial em contextos clínicos preventivos. Além disso, a interpretação dos coeficientes permitiu identificar fatores associados ao risco, como sintomatologia dolorosa, cor da pele, perda de demarcação e endurecimento labial, corroborando evidências já relatadas na literatura odontológica.

No eixo de imagens, o módulo de segmentação do lábio mostrou desempenho satisfatório, assegurando a correta delimitação anatômica e configurando-se como uma base sólida para etapas posteriores. No entanto, o módulo de segmentação de lesões ainda enfrenta desafios significativos relacionados ao desbalanceamento das classes e à baixa confiabilidade das predições, fatores que comprometem sua aplicabilidade imediata. Apesar dessas limitações, a integração entre os dois eixos permanece como uma perspectiva promissora, pois a transformação de achados visuais em variáveis clínicas poderia ampliar a capacidade diagnóstica dos modelos tabulares. Para que isso seja alcançado, será necessário expandir o dataset, empregar técnicas de data augmentation e explorar arquiteturas mais robustas, com validação clínica rigorosa das saídas. Nesse sentido, este trabalho avança ao consolidar a parte tabular e ao estabelecer as bases para o aprimoramento contínuo do eixo imagético, vislumbrando no médio prazo o desenvolvimento de uma ferramenta prática, acessível e de impacto direto no diagnóstico e prevenção do carcinoma de células escamosas em lábios.

Referências

6. REFERÊNCIAS

AL KUWAITI, A. et al. A Review of the Role of Artificial Intelligence in Healthcare. *Journal of Personalized Medicine*, v.13, n.6, p.951, 2023.

ABDOLRASOL, M. G. M. et al. Artificial Neural Networks Based Optimization Techniques: A Review. *Electronics*, v.10, n.21, p.2689, 2021.

BABY, D. et al. Leukocyte classification based on feature selection using extra trees classifier. *Turkish Journal of Electrical Engineering & Computer Sciences*, 2021.

BANSAL, M.; GOYAL, A.; CHOUDHARY, A. A comparative analysis of K-Nearest Neighbor, Genetic, Support Vector Machine, Decision Tree, and Long Short Term Memory algorithms in machine learning. *Decision Analytics Journal*, v.3, n.100071, 2022.

BRITO, L. N. S. et al. Clinical and histopathological study of actinic cheilitis. *Revista de Odontologia da UNESP*, v.48, p.e20190005, 2019.

ESTEVA, A. et al. Dermatologist-level classification of skin cancer with deep neural networks. *Nature*, v.542, n.7639, p.115-118, 2017.

JUBAIR, F. et al. A novel lightweight deep convolutional neural network for early detection of oral cancer. *Oral Diseases*, v.28, n.4, p.1123-1130, 2022.

MEDEIROS, J. R. et al. Clinical and histopathological study of actinic cheilitis. *Revista de Odontologia da UNESP*, v.48, e20190005, 2019.

POITEVIN, A. et al. Actinic cheilitis: proposition and reproducibility of a clinical criterion. *BDJ Open*, v.3, 17016, 2017.

QIU, J. et al. A survey of machine learning for big data processing. *EURASIP Journal on Advances in Signal Processing*, v.67, p.1-16, 2016.

SAITO, T.; REHMSMEIER, M. The Precision-Recall Plot Is More Informative than the ROC Plot When Evaluating Binary Classifiers on Imbalanced Datasets. *PLOS ONE*, v.10, n.3, e0118432, 2015.

STOLTZFUS, J.C. Logistic regression: a brief primer. *Academic Emergency Medicine*, v.18, n.10, p.1099-1104, 2011.

VAN CALSTER, B. et al. Calibration: the Achilles heel of predictive analytics. *BMC Medicine*, v.17, p.230, 2019.

Anexos

Figura 1 – Fluxograma da pipeline tabular

Figura 2 – Importância das variáveis (gráfico de barras)

Figura 3 – Tabela comparativa entre SVC-RBF e Regressão Logística

Figura 4 – Matrizes de confusão (SVM inicial vs Regressão Logística)

Figura 5 – Fluxograma completo (tabular + imagens)

Figura 6 – Frequência das lesões (tabela)

Figura 7 – Curva ROC

Figura 8 – Curva Precision-Recall

CÓDIGO: SB0885

AUTOR: ISAAC MEDEIROS DO AMARAL COSTA

ORIENTADOR: PATRICK CESAR ALVES TERREMATTE

TÍTULO: Modelos de perturbação para auxiliar a compreensão do desenvolvimento do córtex

Resumo

O desenvolvimento do córtex cerebral é impulsionado pela proliferação de progenitores neurais durante tempos do desenvolvimento embrionário, sendo um modelo essencial para compreender malformações congênitas e condições como autismo. Este estudo visa prever os efeitos da perturbação de fatores de transcrição críticos durante o desenvolvimento cortical, gerando um mapa temporal do impacto da supressão destes reguladores na formação do córtex embrionário. Foi implementada uma abordagem computacional baseada em dados de transcriptômica de célula única (scRNA-seq) e epigenômica (scATAC-seq), utilizando pipelines do Seurat para processamento inicial, incluindo filtragem de células e genes de baixa qualidade, normalização, correção de efeitos de batch através do método Harmony e anotação celular baseada em marcadores canônicos. Foi concluída com sucesso a etapa de pré-processamento, integração de dados e visualização através do Quarto, resultando em conjuntos padronizados com remoção eficaz de ruídos técnicos e anotação preliminar de tipos celulares. A qualidade dos dados processados estabelece base robusta para futuras inferências de redes regulatórias gênicas específicas por tipo celular. A expansão para datasets adicionais e aplicação de modelos de redes regulatórias promete viabilizar a caracterização de novos alvos terapêuticos para controle gênico específico durante o desenvolvimento cortical.

Palavras-chave: desenvolvimento cortical; single-cell; redes regulatórias

TITLE: Disturbance models to aid understanding of cortex development

Abstract

The development of the cerebral cortex is driven by the proliferation of neural progenitors during embryonic development, providing an essential model for understanding congenital malformations and conditions such as autism. This study aims to predict the effects of disrupting critical transcription factors during cortical development, generating a temporal map of the impact of suppressing these regulators on embryonic cortex formation. A computational approach based on single-cell transcriptomics (scRNA-seq) and epigenomics (scATAC-seq) data was implemented, using Seurat pipelines for initial processing, including filtering of low-quality cells and genes, normalization, batch effect correction using the Harmony method, and cell annotation based on canonical markers. The preprocessing, data integration, and visualization steps were successfully completed using Quarto, resulting in standardized datasets with effective removal of technical noise and preliminary annotation of cell types. The quality of the processed data establishes a robust basis for future inferences of

cell type-specific gene regulatory networks. Expansion to additional datasets and application of regulatory network models promises to enable the characterization of new therapeutic targets for specific gene control during cortical development.

Keywords: cortical development; single-cell; regulatory networks

Introdução

O desenvolvimento do córtex cerebral tem sido estudado nas últimas décadas, sendo este um excelente modelo para a compreensão de malformações congênitas e algumas outras condições perinatais como autismo. Semelhante aos humanos, quando observados também em camundongos, o processo é impulsionado pela proliferação de progenitores neurais na zona ventricular e subventricular. Os progenitores apicais (APs), localizados na zona ventricular, incluem células-tronco neurais e células gliais radiais, que se dividem para gerar novos progenitores ou neurônios. Desse modo, algumas dessas divisões dão origem aos progenitores intermediários (IPs), que migram para a zona subventricular e sofrem divisões amplificadoras antes de se diferenciarem em neurônios (Di Bella et al., 2021). Esse processo de neurogênese segue um padrão de dentro para fora, formando as camadas corticais progressivamente, enquanto células pós-mitóticas recém-geradas se posicionam acima das mais antigas, moldando a arquitetura do córtex cerebral.

Recentemente, o sequenciamento de célula única revolucionou nossa capacidade de investigar a heterogeneidade celular e os processos de diferenciação durante o desenvolvimento. Esta tecnologia permite a caracterização detalhada dos perfis de expressão gênica em células individuais, revelando trajetórias de diferenciação, estados transitórios e redes regulatórias específicas por tipo celular.

Paralelamente, algumas ferramentas computacionais surgiram para modelar e prever o efeito da perturbação de fatores de transcrição (TFs) em redes reguladoras de genes (Gene Regulatory Networks - GRNs) durante a diferenciação celular (Kinamoto et al., 2023). Essa abordagem combina inferência de redes regulatórias e simulação de perturbações para prever o impacto funcional de fatores de transcrição na decisão do destino celular.

Este trabalho busca prever os efeitos de fatores de transcrição importantes para o desenvolvimento do córtex, tendo assim um mapa temporal do impacto da supressão destes fatores para a formação do córtex durante o período embrionário.

Metodologia

Pré-Processamento e Integração de Dados Single-Cell

Utilizou-se pipelines estabelecidos com Seurat (R) para o processamento inicial dos dados de scRNA-seq e scATAC-seq. O pré-processamento incluiu filtragem de células e genes de baixa qualidade, normalização dos dados de expressão, remoção de efeitos de batch utilizando métodos como Harmony, e anotação preliminar de tipos celulares baseada em marcadores canônicos. Esta etapa estabeleceu a base de dados limpa e padronizada necessária para as análises subsequentes.

Visualização

Utilizando o Quarto foi possível criar uma versão HTML, com recursos e funcionalidades interativos desenvolvidos para facilitar a exploração dos resultados obtidos anteriormente pelas análises no Seurat por outros pesquisadores. Toda a documentação organizada em formato Rmarkdown, garante reprodutibilidade e transparência metodológica.

Resultados e Discussões

Resultados Obtidos

Até o momento, foi concluída com sucesso a primeira etapa metodológica de pré-processamento e integração de dados single-cell. Os dados de scRNA-seq e scATAC-seq no desenvolvimento do córtex foram adequadamente processados e integrados utilizando os pipelines Seurat.

O pré-processamento resultou em conjuntos de dados limpos e padronizados, com remoção eficaz de células e genes de baixa qualidade, normalização adequada dos perfis de expressão e correção satisfatória de efeitos de batch (Figura 2). A anotação preliminar dos tipos celulares foi realizada com base em marcadores canônicos que são observados e classificados na literatura, como é o caso do marcador Pax6, sendo definido como progenitor apical e assim por diante (Figura 1; Yuzwa et al., 2017). Desse modo, fornecendo a base celular necessária para as análises de rede específicas por tipo celular que serão implementadas nas próximas etapas.

Discussão dos Resultados Preliminares

A conclusão bem-sucedida da etapa de pré-processamento estabelece uma base sólida para as análises subsequentes. Assim, a qualidade dos dados integrados que são disponibilizados por laboratórios é crucial para a confiabilidade das redes regulatórias que serão inferidas. Portanto, criar uma padronização adequada entre os diferentes contextos biológicos permitirá comparações robustas e identificação de padrões regulatórios conservados ou específicos de cada condição.

A anotação de tipos celulares representa um passo fundamental, pois a especificidade celular das redes regulatórias é essencial para a identificação de alvos terapêuticos relevantes (Kinamoto et al., 2023). Consequentemente, diferentes tipos celulares apresentam redes regulatórias distintas, e esta heterogeneidade será explorada nas etapas subsequentes para identificar vulnerabilidades específicas que possam ser exploradas em suas nuances.

Conclusão

O projeto, apesar de ainda estar em sua parte inicial, possui dados de expressão de célula única normalizados, assim como essas células e seus respectivos conjuntos sendo anotados usando marcadores clássicos da literatura para cada tipo celular. Finalmente, com o aumento do alcance para outros datasets, é esperado que aplicando o modelo de rede regulatória com uma gama maior de marcadores de célula, seja possível a caracterização de cada vez mais níveis de informação que possam atuar na descoberta do controle de genes específicos que venham causar tipos de distúrbios no processo de desenvolvimento cerebral.

Referências

Di Bella, D.J., Habibi, E., Stickels, R.R. et al. Molecular logic of cellular diversification in the mouse cerebral cortex. *Nature* 595, 554–559 (2021).

Kamimoto, K., Stringa, B., Hoffmann, C.M. et al. Dissecting cell identity via network inference and in silico gene perturbation. *Nature* 614, 742–751 (2023).

Yuzwa, Scott A. et al. Developmental Emergence of Adult Neural Stem Cells as Revealed by Single-Cell Transcriptional Profiling *Cell Reports*, Volume 21, Issue 13, 3970 - 3986 (2017)

Anexos

Figura 1

Figura 2

CÓDIGO: SB1258

AUTOR: EDSON DE CARVALHO FERREIRA FILHO

ORIENTADOR: TETSU SAKAMOTO

TÍTULO: Modelo de aprendizagem de máquina para identificação sorológica do gênero *Leptospira* (Leptospiraceae, Bacteria)

Resumo

O avanço das tecnologias de sequenciamento permitiu identificar padrões genéticos associados a fenótipos específicos em *Leptospira*. Este projeto visa desenvolver um banco de dados genômico e sorológico e utilizar aprendizado de máquina para prever o sorogrupo com base no locus *rfb*. Foram coletadas 1.866 amostras de bancos públicos, sendo 22 sorogrupos com dados suficientes para treinamento. O modelo SVM para classes alcançou métricas de acurácia próximas de 100%. A abordagem mostra-se promissora para diagnósticos moleculares complementares.

Palavras-chave: *Leptospira*, Aprendizado de máquina, Locus *rfb*, Classificação sorológica

TITLE: Machine Learning Model for Serological Identification of the Genus *Leptospira* (Leptospiraceae, Bacteria)

Abstract

Advances in sequencing technologies have enabled the identification of genetic patterns associated with specific phenotypes in *Leptospira*. This project aims to develop a genomic and serological database and apply machine learning to predict serogroups based on the *rfb* locus. A total of 1.866 samples were collected from public databases, with 22 serogroups suitable for training. The SVM model for classes metrics achieved close to 100% accuracy. The approach appears promising as a complementary molecular diagnostic tool.

Keywords: *Leptospira*, machine learning, locus *rfb*, serological classification

Introdução

1. Introdução

A leptospirose é uma zoonose negligenciada de importância global, causada por bactérias do gênero *Leptospira*. Sua ampla diversidade genética e fenotípica, aliada à dificuldade de diagnóstico rápido e preciso, representa um desafio contínuo para a saúde pública (Cerqueira

& Picardeau, 2009). A correta identificação do sorogrupo da cepa infectante é essencial para o monitoramento epidemiológico, o desenvolvimento de vacinas e a escolha de estratégias de controle adequadas. Contudo, os métodos tradicionais baseados em sorotipagem, como o teste de microaglutinação (MAT), apresentam limitações quanto à padronização, tempo de resposta e necessidade de infraestrutura especializada (Ahmed et al., 2015; Gravekamp et al., 1993).

Com os avanços das tecnologias de sequenciamento e da bioinformática, tornou-se possível investigar a base genética da diversidade sorológica em *Leptospira*. Estudos recentes demonstram que a região do genoma conhecida como locus *rfb*, responsável pela síntese de lipopolissacarídeos (LPS), apresenta alta correlação com a classificação sorológica (De la Peña-Moctezuma et al., 1999; Ferreira et al., 2024). Esse locus abriga genes diretamente associados à biossíntese de determinantes antigênicos da parede celular, sendo, portanto, um candidato promissor para o desenvolvimento de marcadores moleculares.

Pesquisas conduzidas por Medeiros et al. (2022) e Ferreira et al. (2024) reforçam a hipótese de que diferentes sorogrupos compartilham perfis genéticos específicos no locus *rfb*, possibilitando a sua utilização em classificações automatizadas por meio de algoritmos computacionais. A aplicação de técnicas de aprendizado de máquina (machine learning) nesse contexto surge como uma alternativa robusta e escalável para a predição sorológica baseada em dados genômicos.

Diante desse cenário, o presente trabalho tem como objetivo geral desenvolver um banco de dados genômico e sorológico curado de amostras do gênero *Leptospira* e aplicar algoritmos de aprendizado de máquina, com foco no locus *rfb*, para predizer o sorogrupo das amostras. Este estudo propõe uma abordagem integrativa entre genômica comparativa e inteligência artificial, visando oferecer uma ferramenta complementar aos métodos diagnósticos tradicionais.

Metodologia

2. Materiais e métodos

Este estudo foi conduzido por meio de uma abordagem computacional integrada, composta por cinco etapas principais: obtenção e curadoria de dados genômicos e sorológicos, identificação de grupos ortólogos relacionados ao locus *rfb*, construção de um banco de dados estruturado, aplicação de algoritmos de aprendizado de máquina e desenvolvimento de uma interface gráfica para uso prático do modelo. Cada etapa foi cuidadosamente planejada para garantir reprodutibilidade, escalabilidade e aplicabilidade dos resultados. A metodologia descrita a seguir reflete o fluxo cronológico e técnico do projeto, desde a fundamentação teórica até a prototipagem funcional de uma ferramenta de predição sorológica baseada em dados genômicos.

2.1 Obtenção de dados

Os dados genômicos e sorológicos utilizados foram obtidos a partir de bancos públicos amplamente utilizados pela comunidade científica. O National Center for Biotechnology Information (NCBI) Reference Sequence Database (RefSeq) (O'LEARY et al., 2016) (<https://www.ncbi.nlm.nih.gov/refseq/>) forneceu 875 genomas de *Leptospira*, enquanto o Bacterial Isolate Genome Sequence Database (Jolley & Maiden, 2010) (BIGSdb,

<http://bigsd.b.pasteur.fr/>) contribuiu com 1.056 amostras adicionais, totalizando 1.931. As sequências foram baixadas em novembro de 2024. Após remoção de duplicatas e curadoria dos metadados sorológicos, restaram 1.357 amostras, sendo 32 sorogrupos identificados. Desses, 22 sorogrupos apresentavam amostras suficientes para o desenvolvimento de classificadores. Dessa forma foram separadas 767 amostras para o treinamento. As sequências foram armazenadas em diretórios organizados por banco de origem (NCBI ou Pasteur) e rotuladas com base em sua classificação taxonômica e sorológica. Sequências de proteínas do locus *rfb* de amostras representativas de cada sorogrupo foram obtidos manualmente a partir de planilhas suplementares disponibilizadas nas publicações de Ferreira et al. (2024)..

2.2 Identificação dos genes do locus *rfb*

Para identificar os genes ortólogos do locus *rfb* entre as amostras, foi utilizado o programa TBLASTN, com comparações entre as sequências proteicas do locus *rfb* das amostras representativas de cada sorogrupo e os genomas completos. Um script em Python foi desenvolvido para automatizar o processo, garantindo reprodutibilidade e escalabilidade. Os resultados, em formato .csv, foram organizados em uma matriz de porcentagem de identidade do melhor hit (corrigida pelo tamanho da query) por gene e por amostra.

2.3 Construção do banco de dados

Com base nos resultados da etapa anterior, foi construída uma base de dados tabular contendo os atributos genômicos derivados do locus *rfb* e os rótulos sorológicos associados. A curadoria incluiu padronização de nomes, tratamento de dados ausentes e exclusão de amostras incompletas. O banco de dados final foi estruturado no formato .tsv, facilitando sua leitura por bibliotecas como pandas e posterior aplicação de algoritmos de aprendizado de máquina. A organização dos dados seguiu os princípios de reprodutibilidade e interoperabilidade, com versionamento dos arquivos intermediários.

2.4 Criação, treinamento e análise do modelo de aprendizagem de máquina

Para as tarefas de predição, foi utilizado o algoritmo Support Vector Machine (SVM) com kernel linear, implementado na biblioteca scikit-learn. O conjunto de dados foi dividido em treino (70%) e teste (30%) com random state = 42. A métrica de avaliação principal foi a acurácia, complementada por F1-score e matriz de confusão. Após validação, o modelo foi salvo com joblib para utilização em etapas posteriores da aplicação. Embora outros algoritmos tenham sido considerados (como Random Forest e KNN), o SVM com kernel linear foi selecionado por seu bom desempenho em conjuntos de dados com dimensionalidade média e número limitado de amostras. Para controle de sobreajuste, foi realizada validação cruzada estratificada com k=5 folds. A análise da matriz de confusão revelou que os erros de classificação se concentraram em sorogrupos com baixa representatividade, indicando a necessidade futura de reequilíbrio das classes ou aplicação de técnicas de oversampling.

2.5 Criação da Interface Gráfica

Uma plataforma web foi desenvolvida para disponibilizar o modelo desenvolvido neste projeto para a comunidade científica. Para isso, foi utilizada Gradio, uma biblioteca python, otimizada para criação de aplicações web com uso de machine learning.

Resultados e Discussões

3. Resultados

3.1 Dados genômicos e sorológicos de *Leptospira* utilizados no trabalho

Durante a análise dos bancos de dados, foram obtidos dados sobre a distribuição sorológica das amostras atualmente disponíveis. Essas amostras foram organizadas em sorogrupos e posteriormente agrupadas em classes (Tabela 1). Além disso, a distribuição de cada espécie por grupos é apresentada na Figura 3 e detalhada na Tabela 2.

A análise dos dados revelou que o grupo mais frequente no banco de dados foi o P1, enquanto o menos frequente foi o S2. As espécies com maior ocorrência foram *Leptospira borgpetersenii*, *L. interrogans* e *L. kirschneri*, todas pertencentes ao grupo P1. Observou-se ainda a presença de diversas espécies pouco frequentes, registradas apenas uma vez, como *L. ognonensis*, *L. bouyouniensis*, *L. semungkisensis* e *L. tipperaryensis*.

No que se refere à classificação por classes de sorogrupo, a mais frequente foi a Classe 1, enquanto a menos frequente foi a Classe 4.

Tabela 3 apresenta a distribuição de diferentes sorogrupos em dez espécies patogênicas de *Leptospira* (*L. interrogans*, *L. borgpetersenii*, *L. santarosai*, *L. kirschneri*, *L. noguchii*, *L. weilii*, *L. yasudae*, *L. kmetyi*, *L. alexanderi* e *L. mayottensis*). Observa-se uma ampla variabilidade na prevalência de sorogrupos entre essas espécies. Por exemplo, *L. interrogans* e *L. borgpetersenii* apresentaram as maiores somas totais de sorogrupos (639 e 324, respectivamente), indicando maior diversidade ou ocorrência de variantes sorológicas nessas espécies. Em contraste, espécies como *L. yasudae*, *L. kmetyi* e *L. alexanderi* exibiram contagens reduzidas (8, 10 e 6, respectivamente).

Alguns sorogrupos, como “Mini” e “Tarassovi”, apresentaram ampla distribuição entre espécies, sugerindo potencial ubiquidade ou relevância epidemiológica. Outros, como “Bim”, “Codice”, “Ranarum”, “Semaranga”, “Muenchen”, “Malaysia”, “Lyme”, “Hurstbridge”, “bairam ali”, “Holland” e “Wolffii”, não foram detectados nesta amostragem. A análise evidencia padrões de coocorrência e dominância, fornecendo subsídios para compreender a ecologia e a dinâmica epidemiológica da leptospirose.

Resultados semelhantes foram reportados por Ferreira et al. (2024), que demonstraram que diferentes espécies podem compartilhar sorogrupos, refletindo a plasticidade genética do locus *rfb* e a complexa relação entre classificação sorológica e taxonômica. Assim, a sobreposição entre sorogrupos e espécies reforça a necessidade de abordagens integrativas, combinando dados genômicos e sorológicos para uma caracterização mais precisa.

3.2 Modelo de predição das classes sorológicas de *Leptospira*.

O presente estudo desenvolveu um modelo de aprendizado de máquina para predição de classes de sorogrupo (I, II, III e IV) de *Leptospira* a partir de dados genômicos, utilizando o locus *rfb* como marcador molecular. A análise concentrou-se exclusivamente em amostras do grupo patogênico (P1), uma vez que a classificação em classes atualmente é restrita a esse grupo. A atribuição de classe baseou-se na porcentagem média de identidade de sequência entre os genes do locus *rfb* das amostras analisadas e de representantes previamente classificados. Amostras com identidade $\geq 50\%$ foram associadas à classe correspondente;

aquelas com identidade inferior a 50% foram excluídas da análise. O conjunto final utilizado no treinamento do modelo consistiu em uma matriz de 767 amostras (linhas) e 27 variáveis (genes do locus rfb).

O classificador foi desenvolvido com o algoritmo Support Vector Machine (SVM) com kernel linear, implementado em scikit-learn. O desempenho foi avaliado por validação cruzada estratificada com cinco partições (StratifiedKFold), preservando-se a distribuição original das classes. Cada amostra foi utilizada uma única vez para validação e nunca avaliada no mesmo subconjunto em que foi treinada, reduzindo o viés de avaliação. A métrica F1-score ponderada atingiu 0,9973, com precisão de 0,9974 e sensibilidade de 0,9973. A matriz de confusão agregada (Figura 6) apresentou valores predominantemente na diagonal principal, com apenas dois erros envolvendo confusão entre as classes 2 e 3, enquanto as classes 1 e 4 obtiveram desempenho perfeito.

A segunda etapa, referente à predição de sorogrupos, encontra-se em desenvolvimento e não foi avaliada nesta fase. Também foi possível explorar a distribuição das amostras sem identificação, que se concentraram principalmente nas classes 3 e 4 (Figura 4), enquanto amostras externas ao conjunto de treinamento mostraram distribuição mais homogênea (Figura 5).

3.3 Interface web

O presente estudo teve, entre seus objetivos, o desenvolvimento de uma interface web que fosse simultaneamente eficiente e de fácil utilização. Ao término do projeto, foi implementada uma interface adaptável por meio da biblioteca Gradio em Python (ABID et al., 2022), amplamente empregada em aplicações de modelos de aprendizado de máquina.

A ferramenta desenvolvida permite a inserção de arquivos no formato FASTA (.fna ou .fas) e, com a execução de um único comando, realiza as predições correspondentes. Os resultados incluem a porcentagem de confiança associada à predição e a distribuição de probabilidades para cada classe avaliada, proporcionando ao usuário uma análise clara e objetiva.

4. Discussão

As distribuições obtidas a partir do banco de dados evidenciaram uma maior ocorrência de amostras pertencentes ao grupo P1 e à Classe 1, indicando predominância de amostras patogênicas. Observou-se também que a maior parte das amostras disponíveis permanece não classificada, tanto em nível de classe quanto de sorogrupo.

Os resultados obtidos evidenciam o potencial da abordagem genômica para a identificação de sorogrupos de *Leptospira*. Comparativamente ao MAT (Microscopic Agglutination Test), considerado o padrão-ouro, a classificação baseada em aprendizado de máquina mostrou-se mais rápida, objetiva e reprodutível. Enquanto o MAT exige infraestrutura laboratorial, cultivo bacteriano e expertise técnica, o modelo proposto opera a partir de dados genômicos brutos, reduzindo tempo e custos (Levett, 2001; Picardeau, 2017).

A escolha do locus rfb mostrou-se apropriada, pois este já havia sido associado à determinação antigênica em *Leptospira* (De la Peña-Moctezuma et al., 1999; Ferreira et al., 2024). Os erros observados sugerem limitações decorrentes do teste, ou na etapa de treinamento do modelo em questão.

Apesar de ainda preliminares, os resultados apontam que o método pode complementar técnicas sorológicas, especialmente em contextos de vigilância epidemiológica e desenvolvimento vacinal. A ampliação e mudança na configuração do conjunto de treino, além da integração de dados de diferentes fontes (genéticos, fenotípicos e clínicos), poderá refinar a acurácia, assim como o uso de técnicas de balanceamento (oversampling) e algoritmos alternativos, como Random Forest ou deep learning.

Além disso, a metodologia proposta pode ser extrapolada para outras bactérias patogênicas que apresentam variação antigênica dependente de loci de superfície, como *Escherichia coli* (antígenos O e K), *Salmonella enterica* (antígenos O) e *Vibrio cholerae* (antígenos O), ampliando seu impacto no campo da genômica aplicada à saúde pública.

Conclusão

5. Conclusão

Este estudo apresentou uma estratégia computacional baseada em aprendizado de máquina para a classificação de *Leptospira* a partir do locus *rfb*. Embora o classificador tenha alcançado desempenho elevado em métricas de treino (acurácia média 88%), sua performance em amostras externas foi limitada (40% de acertos em 10 amostras), ressaltando a necessidade de validação adicional e ajustes metodológicos.

Ainda assim, a construção de um banco de dados curado, associada a um protótipo funcional em interface web, estabelece um marco inicial para o desenvolvimento de ferramentas diagnósticas e epidemiológicas. Essa abordagem tem potencial de ser estendida a outras espécies bacterianas com plasticidade genômica e variabilidade em loci determinantes de superfície, representando um avanço no uso da bioinformática aplicada à vigilância e ao controle de doenças infecciosas.

Referências

6. Referências

ABID, A.; ALFOQAHAA, S.; ZHANG, C.; ZHENG, J.; ZHANG, M.; ALAMA, M.; YALA, A.; ZHANG, J. Gradio: Hassle-Free Sharing and Testing of ML Models in the Wild. *Journal of Machine Learning Research*, v. 23, p. 1-6, 2022. Disponível em: <<https://www.jmlr.org/papers/v23/21-0631.html>>. Acesso em: 29 ago. 2025.

ALTSCHUL, S. F.; GISH, W.; MILLER, W.; MYERS, E. W.; LIPMAN, D. J. Basic local alignment search tool. *Journal of Molecular Biology*, v. 215, n. 3, p. 403–410, 5 out. 1990. DOI: 10.1016/S0022-2836(05)80360-2.

Picardeau, M. (2017). Virulence of the zoonotic agent of leptospirosis: still terra incognita? *Nature Reviews Microbiology*, 15(5), 297–307. <https://doi.org/10.1038/nrmicro.2017.5>

ABELA-RIDDER, Bernadette; SIKKEMA, Reina; HARTSKEERL, Rudy A. Estimating the burden of human leptospirosis. *International Journal of Antimicrobial Agents*, PMID: 20688484, v. 36 Suppl 1, p. S5-7, nov. 2010.

AHMED, Ahmed; FERREIRA, Ana S.; HARTSKEERL, Rudy A. Multilocus sequence typing (MLST): markers for the traceability of pathogenic *Leptospira* strains. *Methods in Molecular Biology* (Clifton, N.J.), PMID: 25399108, v. 1247, p. 349–359, 2015.

ANDREWS, Simon. FastQC: a quality control tool for high throughput sequence data. Available online at: <http://www.bioinformatics.babraham.ac.uk/projects/fastqc>, 2010.

ASHKENAZY, Haim et al. FastML: a web server for probabilistic reconstruction of ancestral sequences. *Nucleic Acids Research*, PMID: 22661579, v. 40, n. Web Server issue, p. W580–584, jul. 2012.

BANKEVICH, A., NURK, S., ANTIPOV, D., GUREVICH, A. A., DVORKIN, M., KULIKOV A. S., et al. SPAdes: A New Genome Assembly Algorithm and Its Applications to Single-Cell Sequencing. *Journal of Computational Biology*, PMID: 22506599. v. 19, p. 455–477, maio 2012.

CAPELLA-GUTIÉRREZ, Salvador; SILLA-MARTÍNEZ, José M.; GABALDÓN, Toni. trimAl: a tool for automated alignment trimming in large-scale phylogenetic analyses. *Bioinformatics* (Oxford, England), PMID: 19505945, v. 25, n. 15, p. 1972–1973, 1 ago. 2009.

CERQUEIRA, Gustavo M.; PICARDEAU, Mathieu. A century of *Leptospira* strain typing. *Infection, Genetics and Evolution: Journal of Molecular Epidemiology and Evolutionary Genetics in Infectious Diseases*, PMID: 19540362, v. 9, n. 5, p. 760–768, set. 2009.

CSÜÖS, Miklós. Count: evolutionary analysis of phylogenetic profiles with parsimony and likelihood. *Bioinformatics*, PMID: 20551134, v. 26, n. 15, p. 1910–1912, 1 ago. 2010.

DARRIBA, Diego et al. jModelTest 2: more models, new heuristics and parallel computing. *Nature Methods*, PMID: 22847109, v. 9, n. 8, p. 772, ago. 2012.

DAVID, Lawrence A.; ALM, Eric J. Rapid evolutionary innovation during an Archaean genetic expansion. *Nature*, v. 469, n. 7328, p. 93–96, 6 jan. 2011.

DE LA PEÑA-MOCTEZUMA, A. et al. Comparative analysis of the LPS biosynthetic loci of the genetic subtypes of serovar Hardjo: *Leptospira interrogans* subtype Hardjoprajitno and *Leptospira borgpetersenii* subtype Hardjobovis. *FEMS microbiology letters*, PMID: 10474199, v. 177, n. 2, p. 319–326, 15 ago. 1999.

EDGAR, Robert C. MUSCLE: multiple sequence alignment with high accuracy and high throughput. *Nucleic Acids Research*, PMID: 15034147, v. 32, n. 5, p. 1792–1797, 1 mar. 2004.

FERREIRA, LCA, et al. Genetic structure and diversity of the *rfb* locus of pathogenic species of the genus *Leptospira*. *Life Sci Alliance*. 2024 Mar 21;7(6):e202302478. doi: 10.26508/lsa.202302478. PMID: 38514188; PMCID: PMC10958091.

GRAVEKAMP, C. et al. Detection of seven species of pathogenic leptospires by PCR using two sets of primers. *Journal of General Microbiology*, PMID: 8409911, v. 139, n. 8, p. 1691–1700, ago. 1993.

GUINDON, Stéphane et al. New algorithms and methods to estimate maximum-likelihood phylogenies: assessing the performance of PhyML 3.0. *Systematic Biology*, PMID: 20525638, v. 59, n. 3, p. 307–321, maio 2010.

JOSHI, N. A., FASS, J. N. Sickle: A sliding-window, adaptive, quality-based trimming tool for FastQ files. Available at <https://github.com/najoshi/sickle>, 2011.

KATOH, Kazutaka; STANDLEY, Daron M. MAFFT Multiple Sequence Alignment Software Version 7: Improvements in Performance and Usability. *Molecular Biology and Evolution*, PMID: 23329690, v. 30, n. 4, p. 772–780, abr. 2013.

KRIVENTSEVA, Evgenia V. et al. OrthoDB v8: update of the hierarchical catalog of orthologs and the underlying free software. *Nucleic Acids Research*, PMID: 25428351, v. 43, n. Database issue, p. D250-256, jan. 2015.

LA SCOLA, Bernard et al. Partial rpoB gene sequencing for identification of *Leptospira* species. *FEMS microbiology letters*, PMID: 16978348, v. 263, n. 2, p. 142–147, out. 2006.

LEHMANN, Jason S. et al. Leptospiral pathogenomics. *Pathogens (Basel, Switzerland)*, PMID: 25437801, v. 3, n. 2, p. 280–308, 2014.

MEDEIROS, EJS, et al. Genetic basis underlying the serological affinity of leptospiral serovars from serogroups Sejroe, Mini and Hebdomadis. *Infect Genet Evol.* 2022 Sep;103:105345. doi: 10.1016/j.meegid.2022.105345. Epub 2022 Jul 30. PMID: 35917899, 2022

PRICE, Morgan N. et al. FastTree 2 – Approximately Maximum-Likelihood Trees for Large Alignments. *PLoS ONE*, DOI: 10.1371/journal.pone.0009490, v. 5, n. 3, p. e9490, 2010

RONQUIST, Fredrik et al. MrBayes 3.2: Efficient Bayesian Phylogenetic Inference and Model Choice Across a Large Model Space. *Systematic Biology*, PMID: 22357727, v. 61, n. 3, p. 539–542, maio 2012.

SIEVERS, Fabian; HIGGINS, Desmond G. Clustal Omega, accurate alignment of very large numbers of sequences. *Methods in Molecular Biology (Clifton, N.J.)*, PMID: 24170397, v. 1079, p. 105–116, 2014.

SLACK, Andrew T. et al. Identification of pathogenic *Leptospira* species by conventional or real-time PCR and sequencing of the DNA gyrase subunit B encoding gene. *BMC microbiology*, PMID: 17067399, v. 6, p. 95, 2006.

STAMATAKIS, Alexandros. RAxML-VI-HPC: maximum likelihood-based phylogenetic analyses with thousands of taxa and mixed models. *Bioinformatics*, PMID: 16928733, v. 22, n. 21, p. 2688–2690, 1 nov. 2006.

SUYAMA, Mikita; TORRENTS, David; BORK, Peer. PAL2NAL: robust conversion of protein sequence alignments into the corresponding codon alignments. *Nucleic Acids Research*, PMID: 16845082, v. 34, n. Web Server issue, p. W609-612, 1 jul. 2006.

YANG, Ziheng. PAML 4: Phylogenetic Analysis by Maximum Likelihood. *Molecular Biology and Evolution*, PMID: 17483113, v. 24, n. 8, p. 1586–1591, 1 ago. 2007.

ZERBINO, Daniel R.; BIRNEY, Ewan. Velvet: Algorithms for de novo short read assembly using de Bruijn graphs. *Genome Research*, PMID: 18349386, v. 18, n. 5, p. 821–829, 1 maio 2008.

Anexos

Matriz de Confusão para o classificador de classes.

Distribuição de classes

Distribuição de sorogrupos

Frequência de grupos por espécie (Tabela 2)

Frequência de cada sorogrupo agrupado pelas classes (Tabela 1)

Distribuição de sorogrupos por espécie