



CI-IA saúde
passado · presente · futuro

Anais do I Simpósio do CI-IA Saúde da UFMG

Ana Paula Couto da Silva
Zilma Silveira Nogueira Reis
Cristiane dos Santos Dias
Juliano de Souza Gaspar



Realização



Apoio



Belo Horizonte, MG
2023



Anais do I Simpósio CI-IA Saúde da UFMG

Passado, presente e futuro.

25 de setembro de 2023 - Belo Horizonte - MG



Anais do I Simpósio CI-IA Saúde da UFMG

Ana Paula Couto da Silva

Zilma Silveira Nogueira Reis

Cristiane dos Santos Dias

Juliano de Souza Gaspar



**Belo Horizonte
2023**



Esta obra é disponibilizada nos termos da Licença Creative Commons – Atribuição – Não Comercial – Compartilhamento pela mesma licença 4.0 Internacional. É permitida a reprodução parcial ou total desta obra, desde que citada a fonte.

S612 Simpósio CI-IA Saúde da UFMG (1. : 2023 : Belo Horizonte, MG).
Anais do I Simpósio CI-IA Saúde da UFMG [recurso eletrônico] : passado, presente e futuro. / organizado por Ana Paula Couto da Silva, Zilma Silveira Nogueira Reis, Cristiane dos Santos Dias. Juliano de Souza Gaspar. - Belo Horizonte : Faculdade de Medicina da UFMG, 2023.

52f.

Requisitos do sistema: Adobe Reader

1. Inteligência Artificial. 2. Saúde. 3. Assistência à saúde. I. Silva, Ana Paula Couto da. II. Reis, Zilma Silveira. III. Dias, Cristiane dos Santos. IV. Gaspar, Juliano de Souza. V. Anais do I Simpósio do Centro de Inovação em Inteligência Artificial para a Saúde da UFMG.

NLM: W 26.5

Bibliotecário responsável: Marina Nogueira Ferraz. CRB-6/2194

DOI: [10.5281/zenodo.8403348](https://doi.org/10.5281/zenodo.8403348)



Universidade Federal de Minas Gerais

Reitora

Sandra Regina Goulart Almeida

Vice-Reitor

Alessandro Fernandes Moreira



Centro de Inovação em Inteligência Artificial para a Saúde - CI-IA Saúde

Conselho Técnico científico

Wagner Meira Júnior

Coordenador do Aliança Estratégica e Diretor do CIIA-Saúde - UFMG

Virgílio Augusto Fernandes Almeida

Coordenador Acadêmico do Comitê Executivo - UFMG

Antonio Luiz Pinho Ribeiro

Vice-coordenador Acadêmico do Comitê Executivo - UFMG

Gilberto Medeiros Ribeiro

Coordenador de Transferência de Tecnologia do Comitê Executivo - UFMG

Zilma Silveira Nogueira Reis

Coordenadora de Educação e Difusão do Conhecimento do Comitê Executivo - UFMG

Mauro Martins Teixeira

Representante da UFMG - UFMG

Leonardo Pereira Florêncio

Representante no Conselho Técnico-Científico - UNIMED BH

Silvana Márcia Bruschi Kelles

Representante no Conselho Técnico-Científico - UNIMED BH

Fernando Martin Biscione

Representante no Conselho Técnico-Científico - UNIMED BH

Gerência de projetos

Fabiana Costa Pereira Peixoto

Letícia Santos Neto



Comissão Organizadora

Ana Paula Couto da Silva
Zilma Silveira Nogueira Reis
Cristiane dos Santos Dias
Juliano de Souza Gaspar

Universidade Federal de Minas Gerais - UFMG
Universidade Federal de Minas Gerais - UFMG
Universidade Federal de Minas Gerais - UFMG
Universidade Federal de Minas Gerais - UFMG

Comissão Científica

Presidente da Comissão Científica

Ana Paula Couto da Silva

Universidade Federal de Minas Gerais - UFMG

Cristiane dos Santos Dias
Débora Miranda
Derrick M Oliveira
Elaine Carvalho
Eura Martins Lage
Fernando Martin Biscione
Juliano de Souza Gaspar
Julio G Domingues
Silvana Márcia Bruschi Kelles
Zilma Silveira Nogueira Reis

Universidade Federal de Minas Gerais - UFMG
Universidade Federal de Minas Gerais - UFMG
Universidade Federal de Minas Gerais - UFMG
Universidade Federal de Minas Gerais - UFMG
Universidade Federal de Minas Gerais - UFMG
UNIMED
Universidade Federal de Minas Gerais - UFMG
Universidade Federal de Minas Gerais - UFMG
UNIMED
Universidade Federal de Minas Gerais - UFMG

Editor

Juliano de Souza Gaspar

Universidade Federal de Minas Gerais - UFMG



Editorial

Simpósio CI-IA Saúde da UFMG

O Centro de Inovação em Inteligência Artificial para a Saúde da UFMG (CI-IA Saúde) visa a pesquisa e o desenvolvimento de soluções avançadas de inteligência artificial (IA), capazes de auxiliar profissionais de saúde no diagnóstico e tratamento de doenças, e orientar gestores de saúde na programação de ações de prevenção e organização da assistência à saúde. Este centro foi aprovado no Edital FAPESP: Chamada de Propostas FAPESP – MCTIC - CGI.BR para Centros de Pesquisas Aplicadas em Inteligência Artificial em fevereiro de 2021 e conta, atualmente, com a participação de 130 pesquisadores de diferentes Instituições Brasileiras.

Anualmente o Centro promove um evento com o objetivo principal reunir os pesquisadores, profissionais e alunos, tanto do setor público como privado, para promover um amplo debate de ideias, fundamentos e aplicações relacionados ao uso da Inteligência Artificial na área da Saúde. O evento deste ano contou com 120 participantes e ocorreu no dia 25 de setembro de 2023 no Instituto de Ciências Exatas no Campus Pampulha da Universidade Federal de Minas Gerais.

Ao todo, o evento contou com a submissão e avaliação de 26 resumos científicos (short papers) e as revisões de pares foram realizadas com o apoio de 11 revisores voluntários. Ao final da avaliação, foram selecionados para publicação 22 trabalhos, que foram apresentados ao público no formato de pôsteres digitais. Os trabalhos mais bem avaliados foram convidados para apresentação oral no evento e receberam certificados de menção honrosa.

A programação do evento contou com mesas redondas onde os projetos de pesquisas apoiados pelo Centro foram apresentados e foram debatidos as competências necessárias para a utilização da Inteligência Artificial na área da Saúde, além da palestra proferida pelo professor Ricardo Cruz-Correia, Universidade do Porto, que abordou os desafios e oportunidades da IA generativa para o ensino e pesquisa em Saúde.

Deixo aqui registrado meu agradecimento aos alunos, professores e pesquisadores que dedicam parte do tempo de suas agendas para contribuir voluntariamente para o sucesso deste evento.

Profa. Dra. Ana Paula Couto da Silva

Presidente da Comissão Científica

Simpósio do CI-IA Saúde da UFMG



Sumário

Editorial	5
Oficina: Inteligência Artificial Generativa em Apoio ao Ensino Médico	7
A CNN-based Approach for Assisting Early Mouth Cancer Detection	8
An automatic bidirectional Libras-Portuguese translation system for improving the access of deaf people to medical assistance	10
An Evaluation of Portuguese to Libras Translator Apps Applied to the Health Context	12
Análise bibliométrica das publicações entre 2016 a 2023 sobre o uso da Inteligência Artificial como ferramenta de predição do risco de amputação em pacientes com o pé diabético	14
Aprendizado supervisionado de máquina para classificação automática de Prontuários Médicos da Cardiologia	16
Auditando Triagens Inteligentes: Um método de avaliação para IAs baseadas em questionários	18
Avaliação do desempenho de modelo preditivo de terapia renal substitutiva em pacientes com covid-19.	20
Descoberta de conhecimento em base de dados sobre procedimentos assistenciais em um plano de saúde	22
Desenvolvimento e implementação de sistema de apoio à decisão para melhorar o controle da hipertensão arterial e do diabetes mellitus	24
Detecção, Caracterização Inicial, Caracterização Semântica - DIS: uma nova metodologia para caracterização e análise de drifts temporais em dados de saúde	26
Inteligência Artificial para a Saúde no TikTok: Uma avaliação da informação	28
microRNAs in cerebrospinal fluid associated with alzheimer's disease: a systematic review and pathway analysis using a machine learning approach	30
Perfil de leitos brasileiros em tempos pré, durante e pós pandêmicos	32
Predição de comportamento suicida em crianças e adolescentes em uma instituição de cuidados psiquiátricos	34
Predição de Necessidade de Internação em Terapia Intensiva e Custos de Tratamento de Pacientes com COVID-19	36
Proposição de uma ontologia de domínio para pacientes oncopediátricos	38
Segmentação semântica de ECG por meio de técnicas de self-learning	40
Sign Language Recognition System for Deaf Patients	42
small RNA Metavir: uma ferramenta baseada em aprendizado de máquina para identificação de vírus em amostras complexas independente de análises de similaridade de sequência	44
Somatório de curvas de Kaplan-Meier: um equivalente meta-analítico para sobrevida global que pode se beneficiar significativamente do uso de inteligência artificial.	46
Transformer generativo auto-regressivo para detecção de anomalias em imagens de ressonância magnética de cérebros	48



Oficina: Inteligência Artificial Generativa em Apoio ao Ensino Médico

Como parte das atividades do II Simpósio do CI-IA Saúde da UFMG, 34 docentes da Faculdade de Medicina da UFMG participaram de uma oficina de aperfeiçoamento didático, patrocinada pelo CI-IA Saúde. Os objetivos da atividade foram: a. compreender os conceitos fundamentais da inteligência artificial generativa: o que é a inteligência artificial generativa, como funciona e suas aplicações potenciais no ensino médico; b. identificar casos de uso relevantes no ensino de graduação, na área da medicina: discutir e trocar experiências de forma a reconhecer oportunidades para integrar a IA generativa em atividades de ensino, avaliando benefícios e limitações; e c, avaliar ferramentas e recursos: experimentar, de forma prática, recursos baseados em inteligência artificial generativa para uso em suas atividades de ensino.

O evento ocorreu no dia 26 de setembro de 2023, no Laboratório de Informática da sala 607 da Faculdade de Medicina da UFMG, integrando as atividades do Programa de Desenvolvimento Docente na Faculdade. O Prof. Ricardo João Cruz Correia conduziu as oficinas, transferindo know-how acadêmico acerca do uso do ChatGPT no apoio ao docente, ao discente e administrativos na Faculdade de Medicina da Universidade do Porto, Portugal. Foram três horas de atividades, de 8:00 até 11:00h e, com a participação ativa dos docentes, as tarefas executadas em tempo real foram exemplos de como a inteligência artificial generativa pode ser usada em benefício do docente. Alguns exemplos foram: na proposição de aulas, exercícios, no apoio à correção de exercícios, na análise de artigos científicos e de bases de dados da saúde. A visão crítica sobre o uso do ChatGPT no ensino médico, questões éticas, incertezas, vantagens e desvantagens permeou toda a sessão.

Agradecemos aos apoiadores, Profa. Eliane Gontijo, Profa. Eleonora Druve e à Diretoria da Faculdade de Medicina e ao nosso tutor internacional.

Profa. Dra. Zilma S Nogueira Reis

Diretoria de Educação e Difusão do Conhecimento

II Simpósio do CI-IA Saúde da UFMG



A CNN-based Approach for Assisting Early Mouth Cancer Detection

Maikel M. Rönnau¹, Tatiana W. Lepper², Luara N. Amaral², Pantelis V. Rados², Manuel M. Oliveira¹

¹Instituto de Informática, Universidade Federal do Rio Grande do Sul, RS

²Faculdade de Odontologia, Universidade Federal do Rio Grande do Sul, RS

maikel.ronnau@inf.ufrgs.br, tatiana.lepper@ufrgs.br, luara.amaral@ufrgs.br, pantelis@ufrgs.br, oliveira@inf.ufrgs.br

Abstract. Oral cancer is the sixth most common kind of human cancer. Brush cytology for counting Argyrophilic Nucleolar Organizer Regions (AgNORs) can help early mouth cancer detection, lowering patient mortality. However, the manual counting of AgNORs still in use today is time-consuming and error-prone. We propose a convolutional neural network (CNN) based method to automatically segment individual nuclei and AgNORs in microscope slide images and count the number of AgNORs within each nucleus. Our solution achieved a performance similar to human experts on a set of 291 images from 6 new patients, while significantly speeding up the process.

Keywords: Oral Cancer, Cytopathology, Deep Learning, Image Segmentation, AgNOR Quantification.

1. Introduction

Cytopathology can help to detect early signs of oral cancer development. The number of stained Argyrophilic Nucleolar Organizer Regions (AgNORs) in the nucleus of cells denote how quickly cells are replicating and serves as an indicator of lesions with malignant potential (1). Although the use of immunohistochemistry has gradually surpassed the use of AgNOR staining, comparative studies using proliferative antibody staining (e.g., p53) demonstrate equivalent results between the two methods, in particular in their predictive ability (2). Thus, given its lower cost, AgNOR staining is an attractive option, especially for developing countries.

2. Methodology

2.1 Dataset

We built a dataset for training and evaluating our CNN model consisting of 1,171 images (2,560×1,920-pixel RGB) from 48 patients annotated by specialists. The images correspond to microscope slides containing brushed epithelial cells from the patients' oral mucosa. The cells were collected from borders between normal tissue and abnormal wounds and stained with argyrophilic staining methods (3). This research was approved by the Ethics Committee (certificate number CAAE - 39212420.9.0000.5347).

2.2 Our Nuclei and AgNOR Image Segmentation Model

To arrive at the most appropriate CNN model for our target application, we evaluated 51 network architectures combining seventeen encoders with three decoders and two loss functions. In total, we trained and evaluated 102 models, while exploiting transfer learning for image segmentation (4) to take advantage of large scale pre-trained models. The model consisting of DenseNet-169 + LinkNet with Focal loss performed best, obtaining a Dice score of 0.90 and intersection over union (IoU) of 0.84. The counting of nuclei and AgNORs, performed after segmentation, achieved precision and recall of 0.94 and 0.90 for nuclei, and 0.82 and 0.74 for AgNORs, respectively.

3. Results

We validate the robustness of our solution by comparing its performance against conventional counting (*i.e.*, visual inspection) of nuclei and AgNORs performed by three human experts on a selected set of nuclei in 291 AgNOR stained images from six new patients. The results of this experiment are summarized in Table 1.



Among the experts, the number of identified nuclei ranged from 326 to 377 (agreement of 86.47%), while the number of detected AgNORs ranged from 908 to 1,141 (agreement of 79.58%). Our model identified 361 nuclei and 1,077 AgNORs, achieving an agreement above 90% with each expert. It obtained better agreement with the experts than the agreement among the experts themselves, while being significantly faster.

Table 1. Performance comparison of our model with three human experts on 291 images from 6 new patients. These are 2,560×1,920 RGB images and the time required by our solution was measured on an Nvidia 3090 GPU.

Attribute	Expert 1	Expert 2	Expert 3	Our Solution
Nuclei count	374	326	377	361
AgNORs count	1,017	908	1,141	1,077
Time	≈ 2h	≈ 3h	≈ 2h	2m26s

We use the Intraclass Correlation Coefficient (ICC) to measure the agreement between the experts and our solution, obtaining a calculated ICC of 0.91 for nuclei and 0.81 for AgNORs with 95% confidence intervals of [0.89, 0.93] and [0.77, 0.84], respectively, and p-values < 0.001, confirming its statistical significance. A detailed discussion about the results can be found in (6).

4. Conclusions

We presented an efficient method for automatic segmentation of nuclei and AgNORs and for counting the number of AgNORs inside each nucleus in cytological images. We also introduced an annotated AgNOR-stained image dataset of epithelial cells from the oral mucosa containing 1,171 images from 48 patients (5). To the best of our knowledge, this is the most diverse publicly available AgNOR-stained image dataset in terms of number of patients, and the first for oral cells. On our dataset, our method achieved Dice and IoU scores of 0.90 and 0.84, respectively, indicating very good agreement with the ground truth. On a validation experiment, our solution achieved a performance similar to human experts on a set of 291 images from 6 new patients, while significantly reducing the time required to quantify the number of AgNORs per nuclei. For instance, it can process 100 high-resolution (4.9 megapixel) images under one minute on a laptop equipped with an Nvidia RTX 2060 GPU. As such, it can free experts from this repetitive, time-consuming and error-prone task, and provides a scalable and easily deployed solution. We hope our method and dataset can help practitioners, stimulate new research, and inspire/support a nationwide program in early oral cancer detection.

5. References

1. Jajodia E, Raphael V, Shunyu NB, Ralte S, Pala S, Jitani AK. Brush cytology and AgNOR in the diagnosis of oral squamous cell carcinoma. *Acta cytologica*. 2017 Feb 9;61(1):62-70.
2. Caldeira PC, Aguiar MC, Mesquita RA, do Carmo MA. Oral leukoplakias with different degrees of dysplasia: comparative study of hMLH1, p53, and AgNOR. *Journal of Oral Pathology & Medicine*. 2011 Apr; 40(4):305-11.
3. Trere D. AgNOR staining and quantification. *Micron*. 2000 Apr 1;31(2):127-31.
4. Raghu M, Zhang C, Kleinberg J, Bengio S. Transfusion: Understanding transfer learning for medical imaging. *Advances in neural information processing systems*. 2019;32.
5. Rönnau, M.M., Lepper, T.W., Amaral, L.N., Rados, P.V., Oliveira, M.M., 2023b. Agnor-stained image dataset of epithelial cells from oral mucosa (AgECOM). <https://github.com/maikelroennau/AgECOM>.
6. Rönnau, M.M., Lepper, T.W., Amaral, L.N., Rados, P.V., Oliveira, M.M., 2023. A CNN-based approach for joint segmentation and quantification of nuclei and NORs in AgNOR-stained images. *Computer Methods and Programs in Biomedicine*, 107788. DOI: <https://doi.org/10.1016/j.cmpb.2023.107788>.



An automatic bidirectional Libras-Portuguese translation system for improving the access of deaf people to medical assistance

J. Ramos¹, R. Nepomuceno⁴, L. Fagundes², I. Franco¹, E. Bernardino³, M. A. Murta³, M. Silva⁴, T. L. Gomes⁴, M. S. Marcolino⁵, E. Nascimento¹, R. Prates¹, D. G. Macharet¹, M. F. M. Campos¹

¹Computer Science Department, Universidade Federal de Minas Gerais (UFMG), Belo Horizonte, MG

²Medical School, Universidade Federal de Ouro Preto (UFOP), Ouro Preto, MG

³Faculty of Letters, Universidade Federal de Minas Gerais (UFMG), Belo Horizonte, MG

⁴Computer Science Department, Universidade Federal de Viçosa (UFV), Viçosa, MG

⁵Medical School, Universidade Federal de Minas Gerais (UFMG), Belo Horizonte, MG

jessicalfr@ufmg.br, raphael.nepomuceno@ufv.br, lucca.oliveira@aluno.ufop.edu.br,
isabel.franco@dcc.ufmg.br, elidea@ufmg.br, michellemurta@ufmg.br, michel.m.silva@ufv.br,
thiago.luange@ufv.br, milenamarc@ufmg.br, erickson@dcc.ufmg.br, rprates@dcc.ufmg.br,
doug@dcc.ufmg.br, mario@dcc.ufmg.br

Abstract. This paper presents an overview of the Captar-Libras project—an innovative research and development initiative aimed at enhancing communication between healthcare professionals and deaf patients. The project leverages advanced technologies including computer vision, artificial intelligence, and human-computer interaction to create an automated system to facilitate seamless interaction by translating video streams of individuals using Brazilian Sign Language (Libras) into written Portuguese text, while also enabling the transformation of Portuguese text into video representations using photorealistic human interpreter avatars. A pivotal aspect of this project involves the creation of a groundbreaking dataset that encompasses annotated videos and 3D data, featuring a collection of medical terms in Libras. These terms are specifically curated to align with the discourse often encountered during pre-consultation discussions. The unique challenges inherent to this task are particularly pronounced even in contexts as constrained as healthcare interactions, further underscoring the ambitious nature of the project. As a result of these efforts, the system's development occupies a prominent position at the forefront of advancements in Sign Language Processing. By addressing the intricate nuances of communication within the medical domain, the Captar-Libras project not only underscores the potential for improved healthcare interactions but also contributes substantially to the broader state-of-the-art landscape.

Keywords: Brazilian Sign Language; Libras; Sign Language Translation; Sign Language Production.

Project: Captar-Libras: Sistema de Comunicação por vídeos para surdos aplicado ao pré-atendimento médico (M. Campos)

1. CONTEXT AND MOTIVATION

Deaf individuals face barriers in medical services, which can significantly impact their care. This fact contributes to the greater number of acute adverse events and chronic complications observed in these individuals [1]. In very few services, there might be a human interpreter, which from one hand would mitigate the communication problem but on the other hand brings other relevant issues such as privacy.

1.1 Objectives

To develop an advanced system based on techniques from computer vision, artificial intelligence, and human-computer interaction to provide bidirectional communication between deaf patients and health professionals.

2. WORK DEVELOPMENT

The project's main activities are the creation of the dataset and the development of the computational framework. The Libras video dataset encompasses sentences essential for doctor-patient communication in emergency care. A multi camera recording setup has been assembled, enabling the capturing of



synchronized, high resolution images from three distinct points of view: frontal, lateral, and overhead. Two cameras record the frontal view: one is used to exclusively acquire images of the face of the individual signing in Libras in order to capture facial expression details. Besides video, the setup includes 3D capturing devices aiming at reducing eventual occlusion problems. The recording environment is also equipped with diffuse, uniform lighting and a standardized chromakey background to enable adequate postprocessing.

To construct the dataset, a *corpus* of common questions and answers in medical care was compiled using a Likert-scale questionnaire applied to healthcare professionals. These relevant terms supported the writing of a set of answers for each common question. Videos of interpreters translating those phrases into Libras were recorded to provide a form of mirroring teleprompter for the deaf volunteers who will be recorded for the dataset. The project is currently under review by the UFMG Ethics and Research Committee (CEP/UFMG) and recordings with volunteers should start soon. Concurrently, the bidirectional translation system is being developed using public sign language datasets, and the models will be fine tuned on our dataset.

The translation system comprises two main pipelines, each one responsible for the translation in one direction. The first pipeline has two stages: the first stage is composed of convolutional and recurrent neural networks combined [2] to extract motion features from the video and categorize them into gloss sequences (textual signal labels). In the second stage, this sequence is converted into Portuguese text, employing a transformer architecture [3]. The second pipeline converts Portuguese text into a sequence of signs and generates an avatar of a human interpreter. This pipeline comprises three stages: the first stage translates text into skeletons of the torso, upper limbs and hands of the corresponding signs, the second, parallel with the first, translates text into coherent facial expressions, and the third adequately combines torso, arm and hand skeletons with facial expressions and produces a photorealistic human interpreter avatar [4].

The system, which will run on a computer equipped with a GPU (Graphics Processing Unit), also features an interface similar to videoconferencing systems, displaying images of the connected individuals and photorealistic human interpreter avatar, while maintaining a text record of the translations - both text and video - performed during a session.

3. CHALLENGES AND LEARNINGS

A pivotal challenge is to establish a comprehensive dataset for a prototype enabling functional bidirectional translation. The *corpus* creation engaged healthcare professionals and Libras experts, involving multiple iterations among them. Furthermore, Sign Language Processing (SLP) represents an emerging field within artificial intelligence, with existing systems still displaying low performance when compared to those handling written language translation. The majority of work in this area is quite recent, with few addressing continuous translation or sign language production. Developing such a system, thus, resides at the forefront of the state-of-the-art in SLP.

4. REFERENCES

1. Bartlett G, Blais R, Tamblyn R, Clermont RJ, MacGibbon B. Impact of patient communication problems on the risk of preventable adverse events in acute care settings. *Canadian Medical Association Journal (CMAJ)*. 2008 Jun 3;178(12):1555-62.
2. Min Y, Hao A, Chai X, Chen X. Visual Alignment Constraint for Continuous Sign Language Recognition. *International Conference on Computer Vision (ICCV)*, 2021.
3. Yin K, Read J. Better Sign Language Translation with STMC-Transformer. *International Conference on Computational Linguistics. International Committee on Computational Linguistics*; 2020.
4. Chan C, Ginosar S, Zhou T, Efros A. Everybody Dance Now. *International Conference on Computer Vision (ICCV)*, 2019.



An Evaluation of Portuguese to Libras Translator Apps Applied to the Health Context

Julia Manuela G. Soares (aluna)¹, Isabel F. de Carvalho (aluna)¹, Raquel O. Prates (orientadora)¹,
Elidéa L. A. Bernardino (co-orientadora)², Milena S. Marcolino (co-orientadora)³

¹Departamento de Ciências da Computação, UFMG, Belo Horizonte, MG

²Faculdade de Letras, UFMG, Belo Horizonte, MG

³Faculdade de Medicina, UFMG, Belo Horizonte, MG

juliamanu10@gmail.com, isabelfc95@gmail.com, rprates@dcc.ufmg.br, elideabernardino@gmail.com,
milenamarc@gmail.com

Resumo. Esta pesquisa avalia a eficácia de aplicativos de tradução automática na conversão do português para a Língua Brasileira de Sinais (Libras) no contexto da saúde, especificamente no pré-atendimento médico. Estudos existentes destacam barreiras significativas de comunicação entre indivíduos surdos e profissionais de saúde no Brasil, o que pode impactar negativamente os serviços médicos, tratamento, prevenção de doenças e até mesmo colocar vidas em perigo. Tecnologias assistivas têm surgido como uma solução potencial para aliviar esses obstáculos. Nosso estudo se concentra em três aplicativos populares de tradução - HandTalk, Rybená e VLibras - explorando sua aplicabilidade nesse cenário. Realizamos duas avaliações desses aplicativos. Primeiramente, empregando o Método de Inspeção Semiótica, identificamos problemas de comunicabilidade. Posteriormente, utilizamos uma combinação de modelos para avaliar a qualidade das traduções. Para gerar essas traduções, utilizamos um corpus criado por especialistas da área médica da nossa equipe de pesquisa, contendo 48 perguntas que podem ocorrer durante o atendimento médico. Os resultados revelam que todos os aplicativos facilitam a comunicação unidirecional e, embora capazes de realizar traduções, apresentaram problemas que comprometem a tradução, podendo ter impactos graves para o atendimento de saúde.

Palavras-chave: tradução automática, língua brasileira de sinais (Libras), surdo, acessibilidade, saúde.

Abstract. This research assesses the effectiveness of automatic translation applications in converting Portuguese into Brazilian Sign Language (Libras) within the medical field. Existing studies highlight significant communication barriers between deaf individuals and healthcare providers in Brazil, which can detrimentally impact medical services, treatment, disease prevention, and even endanger lives. Assistive technologies have emerged as a potential solution to alleviate these obstacles. Our study focuses on three popular translation apps - HandTalk, Rybená, and VLibras - exploring their applicability in this scenario. We conducted dual evaluations of these apps. Firstly, employing the Semiotic Inspection Method, we identified communicative issues. Subsequently, we employed a combination of models to assess translation accuracy. We utilized a specialized corpus of 48 medically-relevant questions, crafted by our research team of medical experts. Findings reveal that all apps facilitate one-way communication and, while capable of performing translations, all of them presented problems that compromised translation and could have serious impacts on health care.

Keywords: automatic translation, Brazilian sign language (Libras), deaf, accessibility, health.

Nome do projeto: Captar-Libras: Sistema de Comunicação por vídeos para surdos aplicado ao pré-atendimento médico.

1. CONTEXTO E MOTIVAÇÃO

Diante da considerável população de 2,3 milhões de indivíduos com deficiência auditiva no Brasil, de acordo com informações extraídas da Pesquisa Nacional de Saúde de 2019, emerge uma crescente inquietação em relação às barreiras comunicativas entre pessoas surdas e ouvintes, especialmente em circunstâncias críticas. Dois dos problemas mais proeminentes consistem na ausência de intérpretes e na falta de preparo dos profissionais de saúde. Tais consequências resultam em atendimentos deficientes, com potencial para impactar negativamente a qualidade de vida das pessoas surdas, além de comprometer sua autonomia,



independência e privacidade. Neste cenário, surgem as tecnologias assistivas como forma de mitigar este problema (1), e dentre as mesmas, destacam-se os tradutores automáticos para a Língua Brasileira de Sinais.

1.1 Objetivo

Através de duas análises distintas, o objetivo deste trabalho consiste em verificar e analisar as possíveis limitações dos sistemas de tradução português-Libras, para a área da saúde.

2. ATIVIDADES PRINCIPAIS

Para este estudo, foram analisados três aplicativos: HandTalk, Rybená e VLibras. Inicialmente, foi conduzida uma avaliação para traçar problemas de comunicabilidade, aplicando-se o Método da Inspeção Semiótica. Em seguida, um corpus contendo 48 questões (gerado por especialistas médicos da equipe de pesquisa) que podem ocorrer durante o atendimento médico foi utilizado para a análise das traduções. Dessa forma, as frases foram inseridas em cada um dos aplicativos e classificadas de acordo com seus respectivos erros de tradução: menores, maiores e críticos, através de um modelo de avaliação baseado em pontuação.

3. DESENVOLVIMENTO DO TRABALHO (METODOLOGIA E RESULTADOS)

Em preparação para o trabalho, três aplicativos Android que traduzem Português para Libras foram selecionados após uma pesquisa bibliográfica abrangente. A escolha se baseou na predominância do Android na América do Sul, optando por aplicativos focados na tradução. Os selecionados foram HandTalk, Rybená e VLibras. A Análise envolveu dois estágios. No primeiro, aplicou-se o Método da Inspeção Semiótica (MIS) (2) aos apps, visando avaliar a comunicabilidade dos sistemas e problemas de interação. A segunda análise focou nas traduções geradas pelos aplicativos das questões do corpus de atendimento médico, usando critérios de erro de tradução (3), os quais definimos como alteração no sentido do texto-fonte produzido pela tradução. A avaliação final resultou em categorias para as frases traduzidas, de "sem problemas" a "problemas graves", usando um sistema de pontos (4). A perda de pontos ocorreu com base na gravidade dos erros, com limites máximos para cada tipo. Como resultado, identificamos problemas de tradução em todas as frases, em todos os aplicativos, com o Rybená apresentando mais erros graves.

4. DESAFIOS E APRENDIZADOS

Embora os sistemas de tradução automática Português-Libras atenuem o problema e apresentem certa facilidade de uso, a realização de apenas um lado da tradução (português-Libras) compromete a comunicação, mantendo-a unilateral. Em relação à qualidade da tradução, todos os aplicativos apresentaram problemas que comprometem a tradução, o que pode causar problemas graves, principalmente no contexto de saúde. Além disso, os tradutores recorrem à datilologia para suprir sinais ausentes e têm problemas com a marcação de interrogações. Essas observações realçam avanços e a necessidade de melhorias contínuas.

5. REFERÊNCIAS

1. Ossada SAR, Ossada KT, Ossada JJC Jr, Issa B. A colaboração de Software para auxiliar na comunicação de surdos em hospitais. Revista Brasileira em Tecnologia da Informação. 2021;3(1):2-13.: <https://www.fateccampinas.com.br/rbti/index.php/fatec/article/view/56>
2. de Souza CS, Leitão CF, Prates RO, Silva EJ. The Semiotic Inspection Method. In: Proceedings of VII Brazilian Symposium on Human Factors in Computing Systems (Natal, RN, Brazil) (IHC '06). ACM; 2006. p. 148-157. doi:10.1145/1298023.1298044.
3. Reis LS, Araújo TMU, Aguiar YPC, Lima MACB. Evaluating Machine Translation Systems for Brazilian Sign Language in the Treatment of Critical Grammatical Aspects. In: Proceedings of the 19th Brazilian Symposium on Human Factors in Computing Systems (Diamantina, Brazil) (IHC '20). ACM; 2020. p. Article 46. doi:10.1145/3424953.3426536.
4. Teixeira MdS. O Jogo da Avaliação: Um Estudo Prático sobre Tradução Automática. 2018. Available from: <https://doi.org/10.17771/PUCRIO.ACAD.41711>.



Análise bibliométrica das publicações entre 2016 a 2023 sobre o uso da Inteligência Artificial como ferramenta de predição do risco de amputação em pacientes com o pé diabético

H. Barbosa¹, N. Betancourt¹, I. Piassi¹, L. Cisneros²

^{1,2}Universidade Federal de Minas Gerais, Belo Horizonte, MG

hugobarbosa1234@gmail.com, nadiarantes@gmail.com, isabellapiassi@gmail.com, ligialoyola@gmail.com

Resumo: Pacientes com úlceras do pé diabético (PD) têm prognósticos ruins de risco de amputação e mesmo sobrevivência. Prevenir e curar essas lesões é um desafio para a equipe de saúde. A inteligência artificial (IA) como ferramenta tecnológica inovadora pode auxiliar no diagnóstico e predição de desfechos contribuindo para a tomada de decisões no manejo desses pacientes. Estudos bibliométricos fornecem um panorama da produção científica existente, questões de destaque e lacunas na literatura. Nosso objetivo foi realizar uma revisão bibliométrica na plataforma Web Of Science (WoS), sobre a produção científica relacionada ao uso da IA no manejo de pacientes com PD. A busca foi realizada na aba de pesquisa avançada, utilizando o campo “tópicos” com os termos “*artificial intelligence*”, “*machine learning*” e “*diabetic foot*” combinados com os operadores booleanos *OR* e *AND* e limitada aos anos de 2016 a 2023. O software VOSviewer Copyright foi usado para analisar os indicadores bibliométricos. Um total de 97 estudos resultou da busca e indicaram a relevância de autores, periódicos, países, instituições, redes de colaboração, categoria da WoS e palavras-chave sobre o tema e evidenciaram um aumento nos últimos anos do uso, publicações e participações da comunidade científica mundial na abordagem do PD usando IA.

Palavras-chave: Inteligência Artificial; Pé Diabético; Tomada de decisões.

Nome do projeto: “Análise bibliométrica das publicações sobre o uso da Inteligência Artificial como ferramenta de predição do risco de amputação em pacientes com pé diabético”. **Obs:** por se tratar de estudo metodológico que utiliza dados secundários e não envolve seres humanos, não necessita de aprovação ética.

1. CONTEXTO E MOTIVAÇÃO

Pacientes com úlceras do pé diabético (PD) têm prognósticos ruins de risco de amputação e mesmo sobrevivência (1). Prevenir e curar essas lesões é um desafio para a equipe de saúde (2). O uso da inteligência artificial (IA) na área da saúde é recente e ainda são poucos os estudos sobre sua aplicabilidade e impactos na rotina do profissional de saúde, incluindo na assistência a pacientes com PD (3,4). Um levantamento quantitativo da bibliografia permite apontar indicadores de produção científica e direções de estudos. Isso é possível através da realização de uma análise bibliométrica da literatura mundial. Nosso objetivo foi realizar uma revisão bibliométrica na plataforma Web Of Science (WoS), sobre a produção científica relacionada ao uso da IA no manejo de pacientes com PD.

2. ATIVIDADES PRINCIPAIS

Acesso ao portal de periódicos CAPES. Treinamento de uso da WoS e do software VOSviewer Copyright e seus respectivos recursos. Montagem da estratégia de busca e refinamento a partir da identificação das palavras-chave presentes em publicações relevantes sobre o tema e da pesquisa nos Descritores de Ciências da Saúde (DECS/ MeSH) para a correta conversão dos termos de português para o inglês. Análise dos resultados por relevância com o tema, classificados por número de publicações e de citações. Organização dos resultados em gráficos, tabelas e grafos.

3. DESENVOLVIMENTO DO TRABALHO

Estudo metodológico do tipo revisão bibliométrica com dados da plataforma WoS. Foram incluídos trabalhos publicados entre 2016 e 2023. Os termos “*artificial intelligence*”, “*machine learning*”, “*diabetic foot*” foram



utilizados na estratégia de busca, combinados com os operadores booleanos *OR* e *AND*. A busca foi realizada na aba de pesquisa avançada, utilizando o campo “tópicos”. O software VOSviewer Copyright foi usado para analisar os indicadores bibliométricos de autores, periódicos, países, instituições, redes de colaboração, categoria da WoS e palavras-chave. Os dados coletados foram avaliados quanto à sua relevância, classificados por quantidade de publicações e de citações.

Resultados: A busca resultou em um total de 97 estudos, sendo que o maior número de publicações ocorreu no ano de 2022 (n=41). O número de documentos publicados aumentou em 9600% de 2016 (1 documento) até julho de 2023 (96 documentos). Os autores Sawal Hamid Md Ali e Muhammad E.H. Chowdhury foram os que mais publicaram. *Computers in Biology and Medicine* e *Sensors* foram identificados como os periódicos com mais publicações (n=6). A publicação mais citada foi “*Skin-Inspired Antibacterial Conductive Hydrogels for Epidermal Sensors and Diabetic Foot Wound Dressings*” (n=307). Os países com maior quantidade de estudos publicados foram: Estados Unidos, Índia e Inglaterra. As instituições de destaque foram: *Qatar University* e *Universiti Kebangsaan Malaysia*. Os 512 autores da rede de colaboração foram agrupados em 75 *clusters* distintos. Cada um desses *clusters* revela padrões únicos de colaboração e produtividade entre os autores, contribuindo para uma compreensão mais completa da dinâmica dessa rede colaborativa. No gráfico de densidade resultante dessa análise, é evidente a presença de 3 grupos proeminentes. A categoria da WoS “*Engineering Electrical Electronic*” contabilizou a maior quantidade de publicações, totalizando 17 estudos. As palavras-chave mais frequentes foram “*machine learning*” (n=36), “*deep learning*” (n=23) e “*diabetic foot*” (n=21).

4. DESAFIOS E APRENDIZADOS

Realizar uma análise bibliométrica sobre o uso da IA no manejo do PD começou pelo desafio da triagem criteriosa de descritores, suas variações e combinações que levassem aos estudos relacionados a interseção da IA com o manejo de pacientes com PD. Na sequência, os desafios vieram na organização dos dados de forma a espelharem a realidade da literatura disponível na WoS sobre o tema com confiabilidade e clareza. O sucesso do método permitiu identificar autores mais ativos, países e instituições que mais produziram conhecimento e os grupos de colaboração que mais contribuíram para determinada pesquisa, o que é muito útil para leitores que estão iniciando na área. Identificamos também os periódicos que viabilizaram a divulgação pública das pesquisas e os termos (palavras-chave) mais recorrentes nesta literatura científica. A identificação de periódicos com alta contagem de publicações e cocitação na área são informações importantes para pesquisadores que buscam periódicos de alta qualidade para suas publicações sobre o tema. Percebemos com esta análise o movimento crescente do uso da IA no manejo de pacientes com PD e as lacunas de pesquisas para o futuro que possam melhorar a qualidade da assistência a essa população utilizando sistemas inteligentes.

5. REFERÊNCIAS

1. Costa RHR, Cardoso NA, Procópio RJ, Navarro TP, Dardik A, de Loiola Cisneros L. Diabetic foot ulcer carries high amputation and mortality rates, particularly in the presence of advanced age, peripheral artery disease and anemia. *Diabetes Metab Syndr*. 11 Suppl 2: © 2017. Published by Elsevier Ltd.; 2017. p. S583-S7.
2. Toscano C, Sugita T, Rosa M, Pedrosa H, Rosa R, Bahia L. Annual Direct Medical Costs of Diabetic Foot Disease in Brazil: A Cost of Illness Study. *International Journal of Environmental Research and Public Health*. 2018;15(1):89.
3. Tulloch J, Zamani R, Akrami M. Machine Learning in the Prevention, Diagnosis and Management of Diabetic Foot Ulcers: A Systematic Review. *IEEE Access*. 2020;8:198977-9000.
4. de Loiola Cisneros L. Redes Neurais no cuidado de pacientes: ficção ou realidade?/Neural Networks in Patient Care: Fiction or Reality? *Rev Cienc Saude*, 2019. p. 1-2.



Aprendizado supervisionado de máquina para classificação automática de Prontuários Médicos da Cardiologia

G. C. da Silva^{1,2}, S. C. Cazella²

¹ Hospital São Francisco, Porto Alegre, RS

² Programa de Pós-Graduação em Tecnologias da Informação e Gestão em Saúde da Universidade Federal de Ciências da Saúde de Porto Alegre, Porto Alegre, RS

gabriel.constantin@ufcspa.edu.br, silvioc@ufcspa.edu.br

Resumo. A gestão de documentos eletrônicos na área da saúde é desafiadora devido ao grande volume de informações. Nesse contexto, a classificação automática de textos é essencial para lidar com a quantidade crescente de documentos eletrônicos. Este artigo descreve a pesquisa na qual se propõe o desenvolvimento de um modelo baseado em aprendizado supervisionado de máquina para classificar textos de prontuários eletrônicos de pacientes provenientes de um hospital de cardiologia. O estudo seguirá as seguintes etapas: 1) Geração de um dataset sintético composto por 50 amostras com prontuários similares aos prontuários eletrônicos disponíveis em um hospital de cardiologia utilizando a ferramenta ChatGPT; 2) Tratamento dos textos usando técnicas de processamento de linguagem natural (PLN); 3) Aplicação dos prontuários sintéticos pré-processados para classificação usando os algoritmos Perceptron Multicamadas, Árvores de Decisão, K-Nearest Neighbor e Long-Short Term Memory para predição de diagnóstico; 4) Avaliação dos modelos preditivos pelas métricas de acurácia, precisão e sensibilidade. O resultado esperado é um modelo preditivo robusto capaz de classificar prontuários médicos da cardiologia, trazendo benefícios como auxiliar na tomada de decisões clínicas e melhorar a eficiência dos processos médicos.

Palavras-chave: Processamento de linguagem natural; Aprendizado Supervisionado de Máquina; Prontuários Médicos Eletrônicos.

Abstract. Managing electronic documents in the health area is challenging due to the large volume of information. In this context, automatic text classification is essential to deal with the growing amount of electronic documents. This article describes the research proposed to develop a model based on supervised machine learning to classify texts from electronic medical records of patients from a cardiology hospital. The study will follow these steps: 1) Generation of a synthetic dataset composed of 50 samples with medical records similar to the electronic medical records available in a cardiology hospital using the ChatGPT tool; 2) Treatment of texts using natural language processing (NLP); 3) Application of texts from pre-processed synthetic medical records to Multilayer Perceptron, Decision Trees, K-Nearest Neighbor, and Long-Short Term Memory algorithms to generate a predictive diagnostic model; 4) Evaluation of the predictive models through accuracy, precision, and sensitivity metrics. The expected result is a robust predictive model capable of classifying cardiology medical records, bringing benefits such as assisting in clinical decision-making and improving the efficiency of medical processes.

Keywords: Natural language processing; Machine learning; Electronic Medical Records.

Nome do projeto: Aprendizado de máquina supervisionado para classificação automática de Prontuários Médicos da Cardiologia.

1. CONTEXTO E MOTIVAÇÃO

O prontuário eletrônico, preenchido por médicos, enfermeiros e técnicos, desempenha um papel fundamental na obtenção de dados organizados sobre a evolução clínica, contendo dados valiosos para pesquisas retrospectivas (1). Porém, os dados não são estruturados e possuem repetições e redundâncias. Neste contexto, o uso das técnicas de processamento de linguagem natural (PLN) permite o processamento e análise de documentos seguido da extração de informações relevantes (2). Algumas técnicas modernas de PLN têm se destacado, como a extração de palavras-chave, reconhecimento de entidades nomeadas, análise semântica e segmentação de tópicos (3).

1.1 Objetivo



Propor um modelo preditivo de aprendizado supervisionado de máquina para a classificação automática de textos de prontuários médicos eletrônicos, esperando-se alcançar uma alta taxa de acurácia na classificação.

2. ATIVIDADES PRINCIPAIS

- Criação do conjunto de dados sintéticos para simular os dados de prontuários reais.
- Pré-processamento dos textos a fim de retirar elementos do texto não utilizados nos modelos.
- Aplicação e avaliação de diferentes modelos de aprendizado de máquina.

3. DESENVOLVIMENTO DO TRABALHO

Pretende-se utilizar prontuários médicos eletrônicos reais assim que receber aprovação no Comitê de Ética em Pesquisa da Santa Casa de Porto Alegre. Realizaremos o pré-processamento dos dados sintéticos para formatação dos textos para adequá-los para análise. Com os dados pré-processados, serão treinados modelos para classificar os prontuários médicos com base em seu conteúdo. Utilizaremos bibliotecas NLTK (Natural Language Toolkit¹), spaCy² e scikit-learn³ porque são amplamente utilizadas nesta tarefa e possuem uma grande comunidade de usuários. O desempenho dos modelos será avaliado por acurácia, precisão e sensibilidade (4).

4. DESAFIOS E APRENDIZADOS

O desafio enfrentado foi a criação do dataset sintético para simular informações reais, além de configurar o modelo para aprender padrões e realizar classificações precisas. O aprendizado obtido com este trabalho envolveu as etapas de preparação do conjunto de prontuários médicos eletrônicos utilizados nos modelos de classificação, segundo o diagnóstico atribuído.

5. AGRADECIMENTO

A equipe do projeto agradece à Fundação de Amparo à Pesquisa do Rio Grande do Sul pelo apoio financeiro através do Projeto FAPERGS 22/2551-0000390-7 (RITE CIARS).

6. REFERÊNCIAS

1. Peek N, Holmes JH, Sun J. Technical challenges for big data in biomedicine and health: data sources, infrastructure, and analytics. Yearb Med Inform. 2014 Aug 15;9(1):42-7. doi: 10.15265/IY-2014-0018. PMID: 25123720; PMCID: PMC4287098.
2. Meystre SM, Savova GK, Kipper-Schuler KC, Hurdle JF. Extracting information from textual documents in the electronic health record: a review of recent research. Yearb Med Inform. 2008;128-44. PMID: 18660887.
3. Beliga S, Meštrović A, Martinčić-Ipšić S. An Overview of Graph-Based Keyword Extraction Methods and Approaches. J Inform Organ Sci. 2015;39:1-20.
4. Hossin M, Sulaiman MN. A review on evaluation metrics for data classification evaluations. Int J Data Min Knowledge Manage Process. 2015;5(2):1.

¹ <https://www.nltk.org/>

² <https://spacy.io/>

³ <https://scikit-learn.org/stable/index.html>



Auditando Triagens Inteligentes: Um método de avaliação para IAs baseadas em questionários

Victoria Estanislau Ramos Oliveira, Pedro Loures Alzamora, Camila Nicola, Derick Matheus Oliveira, Wagner Meira Jr (orientador)

Departamento de Ciência da Computação, UFMG, Belo Horizonte, MG

victoriaeramosdeoliveira@gmail.com, ploures.alzamora@gmail.com, camilanicolads@gmail.com, derick.matheuss@gmail.com, meira@dcc.ufmg.br

Resumo. Esse Resumo expandido descreve o processo de elaboração de um modelo de triagem de TDAH assistido por inteligência artificial (IA) explicável. O processo consistiu na redução do conjunto de features preditoras, aplicação de modelo de IA explicável e criação de um método de auditoria com especialistas da área, seguido pela sua análise.

Abstract. This expanded summary describes the elaboration process of an ADHD screening model assisted by an explainable artificial intelligence (AI). It consisted in features selection, apply the explainable AI model and an auditing algorithm alongside it. This auditing system was analyzed by psychiatrists to validate the results obtained.

Palavras-chave: Inteligência Artificial; Explicabilidade; Auditabilidade; TDAH.

Nome do projeto: IA para saúde

1. CONTEXTO E MOTIVAÇÃO

O Transtorno de Déficit de Atenção com Hiperatividade (TDAH) é um transtorno do neurodesenvolvimento cujos sintomas variam dentre desatenção, impulsividade e hiperatividade. O transtorno afeta cerca de 5% de crianças e 2,5% de adultos, prejudicando o desempenho acadêmico e profissional dos portadores^[1]. Estima-se que, no Brasil, menos de 20% deles possuam o tratamento adequado^[2]. Dessa forma, para contribuir com o acesso ao diagnóstico, bem como ao tratamento, pode-se automatizar o processo de triagem por meio de classificadores, auxiliando profissionais generalistas na tomada de decisão acerca de encaminhamento para especialistas e, ainda, para iniciar tratamentos de forma mais assertiva. Para que isso seja possível, é preciso que, além da acurácia, o classificador possua fácil interpretação^{[3] [4]}.

1.1 Objetivo

O objetivo desse projeto foi, não só criar um sistema para auxiliar o processo de triagem de TDAH de modo explicável, mas elaborar uma metodologia de auditoria para sistemas de predição de características baseadas em questionários.

2. ATIVIDADES PRINCIPAIS

Dentre as atividades exercidas, podemos salientar 5:

- Revisão bibliográfica
- Tratamento da base de dados de triagem de TDAH
- Aplicação de modelo Explainable Boost Machine (EBM)^[5]
- Desenvolvimento de uma nova interpretação para o EBM
- Desenvolvimento de uma metodologia de auditoria para o modelo de triagem de TDAH elaborado

Passando rapidamente por cada uma das atividades: a revisão bibliográfica foi importante para entender como a inteligência artificial está sendo usada no campo da saúde mental e qual o estado da arte no meio de



auditoria e inteligência artificial explicável. Em seguida, foi necessário aplicar uma série de tratamentos na base de dados como dicotomização das respostas, tratamento de dados faltantes e tratamento de dados sensíveis, para então podermos aplicar o modelo de IA explicável escolhido: o EBM. Finalmente, tendo aplicado o modelo foi necessário criar uma nova explicação que fosse interpretável para o público não voltado para a computação, esse modelo de explicação foi utilizado na auditoria onde especialistas da área de saúde avaliaram o sistema de inteligência artificial.

3. DESENVOLVIMENTO DO TRABALHO

Para o desenvolvimento do modelo EBM, foi utilizada uma base de dados de questionários psicométricos de crianças com suspeita de TDAH. O modelo selecionou atributos com correlações superiores a dois desvios padrões da média. Dado isso, seria preciso avaliar o desempenho do EBM, o que foi possível por meio do sistema de auditoria. Inicialmente, nesse sistema os especialistas escolheram dentre as questões selecionadas aquelas que afetariam o diagnóstico de TDAH. Em seguida, foram avaliados exames em comparação com o diagnóstico fornecido pelo modelo e, visando melhorar a interpretação, um gráfico com os pesos dos atributos considerados pelo modelo foi exibido. Por fim, os especialistas informaram se discordavam dos critérios da máquina e se mudariam seu diagnóstico após visualizar a interpretação fornecida.

Após coletadas as impressões dos especialistas, procuramos avaliar as questões selecionadas por eles de acordo com sua relevância. Para isso seguimos duas alternativas: (i) identificamos um subconjunto das questões que todos concordaram, sendo este formado por questões sempre selecionadas e nunca selecionadas e (ii) utilizamos a Teoria de Resposta ao Item (TRI) com o modelo Rasch para avaliarmos as propriedades dos atributos selecionados pelo modelo prevendo construtos latentes.

4. DESAFIOS E APRENDIZADOS

Esse projeto teve dois grandes desafios: Uma base de dados esparsa e desbalanceada e a formulação de um modelo robusto de auditoria. Esses desafios se mostraram pertinentes durante todo o curso do projeto. Respectivamente, a primeira a seleção de features utilizadas para a classificação de TDAH assim como compromete a sensibilidade do modelo e a última envolve desde o contato com especialistas, a tradução de resultados computacionais para linguagem médico até a decisão de quais métricas seriam relevantes de ser avaliadas nesse caso em que os próprios especialistas poderiam não estar em acordo.

No entanto, desse último desafio surgiu o principal aprendizado do projeto que é a elaboração de um modelo que avalia tanto a seleção de features feita pela máquina, a concordância entre os especialistas assim como a relevância das features em relação às escolhas do especialista.

5. REFERÊNCIAS

1. Association, A. P. et al. "DSM-5: Manual diagnóstico e estatístico de transtornos mentais." Artmed Editora (2014).
2. Mattos, P., Rohde, L. A., and Polanczyk, G. V. "Adhd is undertreated in brazil." Brazilian Journal of Psychiatry vol. 34, pp. 513–514 (2012).
3. Lipton, Z. C. "The mythos of model interpretability" CoRR vol. abs/1606.03490 (2016).
4. Lipton, Z. C. "The doctor just won't accept that!" arXiv preprint arXiv:1711.08037 (2017).
5. Nori H, Jenkins S, Koch P, Caruana R. Interpretml: A unified framework for machine learning interpretability. arXiv preprint arXiv:1909.09223. 2019 Sep 19.



Avaliação do desempenho de modelo preditivo de terapia renal substitutiva em pacientes com covid-19.

V.G.J. Ventura (aluna), C.M.V. Andrade (aluno), M.S. Marcolino (orientadora), M.A. Gonçalves (orientador)
Universidade Federal de Minas Gerais, Belo Horizonte, MG
nessachese@yahoo.com.br, claudiovaliense@gmail.com, mgoncalv@dcc.ufmg.br, milenamarc@gmail.com

Resumo. O estudo avaliou modelo preditivo de terapia renal substitutiva (*TRS*) durante internação em 3.030 pacientes adultos com covid-19 admitidos em unidade de terapia intensiva em 25 hospitais, utilizando *Random Forest*. A idade mediana dos pacientes foi de 61 anos [intervalo interquartil 51-71] e 26,7% necessitaram de *TRS*. O desempenho do modelo considerando o conjunto de classes foi AUROC 0,66 (IC95% 0,64-0,68) e Macro-F1 0,56 (IC95% 0,55-0,57). Na avaliação de cada classe individualmente, foi observada *Precision* de 0,51 (IC95% 0,47-0,55), *Recall* de 0,20 (IC95% 0,18-0,22) e *F1* de 0,29 (IC95% 0,28-0,30) para o grupo que necessitou de *TRS*; e *Precision* de 0,76 (IC95% 0,75-0,77), *Recall* de 0,93 (IC95% 0,92-0,94) e *F1* de 0,84 (IC95% 0,83-0,85) para o grupo que não necessitou de *TRS*. Os resultados demonstram que *Recall* e *Precision* são fundamentais para uma compreensão adequada do desempenho do modelo preditivo quando há desequilíbrio entre as classes.

Abstract. The study evaluated a predictive model of kidney replacement therapy (*KRT*) during hospitalization in 3,030 adult patients with COVID-19 admitted to the intensive care unit in 25 hospitals, using *Random Forest*. The median age of patients was 61 years [interquartile range 51-71] and 26.7% required *KRT*. The performance of the model considering the set of classes was AUROC 0.66 (CI95% 0.64-0.68) and Macro-F1 0.56 (CI95% 0.55-0.57). In the evaluation of each class individually, were observed *Precision* of 0.51 (CI95% 0.47-0.55), *Recall* of 0.20 (CI95% 0.18-0.22), and *F1* of 0.29 (CI95% 0.28-0.30) for the group with *KRT* requirement; and *Precision* of 0.76 (CI95% 0.75-0.77), *Recall* of 0.93 (CI95% 0.92-0.94), and *F1* of 0.84 (CI95% 0.83-0.85) for the group without *KRT* requirement. The results demonstrate that *Recall* and *Precision* are fundamental for a proper understanding of the performance of the predictive model when there is an imbalance between the classes.

Palavras-chave: covid-19; terapia renal substitutiva; modelo preditivo.

Nome do projeto: Predição de Desfechos Clínicos e Econômicos por Meio de Representações Semânticas MultiModais de Pacientes Resilientes a *Drifts* Temporais (Marcos André Gonçalves).

1. CONTEXTO E MOTIVAÇÃO

A lesão renal aguda é uma complicação comum em pacientes com covid-19 grave, podendo levar à necessidade de terapia renal substitutiva (*TRS*) [1]. Modelos preditivos podem contribuir para a identificação precoce de pacientes em risco, apoiando decisões clínicas. Franca et al. (2023) desenvolveram um modelo preditivo para *TRS* utilizando *Random Forest* e observaram bom desempenho, avaliado pela área sob a curva característica de operação do receptor (AUROC) [2]. Entretanto, os dados mostraram classes desbalanceadas e AUROC pode não ser a melhor métrica de avaliação nesse contexto. O objetivo deste estudo foi avaliar o modelo de Franca et al. (2023) na predição de *TRS* em base de dados de coorte multicêntrica de hospitais brasileiros.

2. ATIVIDADES PRINCIPAIS

Revisão da literatura, análise e interpretação dos dados, redação do texto científico.

3. DESENVOLVIMENTO DO TRABALHO

Estudo de coorte retrospectiva, que incluiu pacientes adultos consecutivos (≥ 18 anos), com covid-19 confirmada laboratorialmente, internados em unidade de terapia intensiva (UTI) em 25 hospitais, de março de 2021 a agosto de 2022. Gestantes, pacientes transferidos de outro hospital ou para outro hospital que não



fazia parte do estudo, admitidos por outras condições de saúde que não covid-19, em tratamento paliativo, com história de *TRS* prévio ou em *TRS* à admissão foram excluídos. O desfecho primário foi a necessidade de *TRS* durante a internação. As variáveis independentes foram idade, sexo, hipertensão, diabetes, doença renal crônica, índice de fragilidade modificado, escala de coma de Glasgow, creatinina, ureia, plaquetas, uso de aminos, suporte ventilatório não invasivo e invasivo à admissão na UTI, conforme estudo de Franca et al. (2023) [2]. Para esta análise, foi empregado *5-fold stratified cross-validation*. O conjunto de dados foi dividido em 5 partes, mantendo as proporções de classe (*folds*). O processo foi repetido 5 vezes. Em cada iteração, um dos *folds* foi selecionado como conjunto de teste, enquanto os outros 4 foram utilizados para compor o conjunto de treinamento e validação. O treinamento foi utilizado para aprendizado do modelo, enquanto a validação foi empregada para ajuste de parâmetros. Assim como no modelo de Franca et al. (2023), foi aplicado o classificador *Random Forest* [2].

A avaliação de cada classe individualmente e do conjunto de classes foram usadas como métricas. Foram calculadas no conjunto de teste: *Precision* do algoritmo (número de acertos em relação ao número total de vezes que o algoritmo reconheceu a classe específica), *Recall* (o quanto o algoritmo conseguiu identificar corretamente a classe alvo), métrica *F1* (média harmônica entre *Precision* e *Recall*), Macro-F1 e AUROC, com seus respectivos intervalos de confiança (IC).

O estudo foi aprovado pela Comissão Nacional de Ética em Pesquisa (CAAE 30350820.5.0000.0008).

Foram incluídos 3.030 pacientes (idade mediana 61 anos [intervalo interquartilico 51-71], 56,4% eram homens e 26,7% necessitaram de *TRS*). O desempenho do modelo considerando o conjunto de classes foi AUROC 0,66 (IC95% 0,64-0,68) e Macro-F1 0,56 (IC95% 0,55-0,57). Na avaliação de cada classe individualmente, foi observada *Precision* de 0,51 (IC95% 0,47-0,55), *Recall* de 0,20 (IC95% 0,18-0,22) e *F1* de 0,29 (IC95% 0,28-0,30) para o grupo que necessitou de *TRS*; e *Precision* de 0,76 (IC95% 0,75-0,77), *Recall* de 0,93 (IC95% 0,92-0,94) e *F1* de 0,84 (IC95% 0,83-0,85) para o grupo que não necessitou de *TRS*.

Os resultados demonstram que a avaliação de métricas de classes individuais e do grupo são fundamentais para uma compreensão adequada do desempenho do modelo preditivo quando há desequilíbrio entre as classes. A partir das métricas supracitadas, fica mais claro entender o desempenho do modelo, por exemplo, se a predição de necessidade de *TRS* é mais importante, podemos priorizar a *Recall* e *Precision* desta classe, fazendo com que o modelo consiga prever de forma mais eficaz pessoas que necessitam de *TRS*. Na área da saúde, esse conhecimento é essencial, visto que um erro de classificação pode impactar negativamente a evolução dos pacientes, por exemplo, deixar de identificar um paciente com risco de evolução para *TRS* pode influenciar a equipe a não adotar intervenções preventivas e terapêuticas em tempo oportuno.

4. DESAFIOS E APRENDIZADOS

Necessidade de análise metodológica cuidadosa e aplicação de métricas adequadas de avaliação.

5. REFERÊNCIAS

1. Yang L, Li J, Wei W, Yi C, Pu Y, Zhang L, Cui T, Ma L, Zhang J, Koyner J, Zhao Y, Fu P. Kidney health in the COVID-19 pandemic: An umbrella review of meta-analyses and systematic reviews. *Front Public Health*. 2022;10:963667.
2. Franca A, Rocha E, SL Bastos L, Bozza FA, Kurtz P, Maccariello E, et al. Development and Validation of a Machine Learning Model to Predict the Use of Renal Replacement Therapy in 14,374 Patients COVID-19 [Preprint]. SSRN. 2023.
3. Brabec J, Komárek T, Franc V, Machlica L. On model evaluation under non-constant class imbalance. *Comput Sci*. 2020; 12140: 74-87
4. Liu S, Roemer F, Ge Y, Bedrick EJ, Li ZM, Guermazi A, Sharma L, Eaton C, Hochberg MC, Hunter DJ, Nevitt MC, Wirth W, Kent Kwok C, Sun X. Comparison of evaluation metrics of deep learning for imbalanced imaging data in osteoarthritis studies. *Osteoarthritis Cartilage*. 2023 Sep;31(9):1242-1248.



Descoberta de conhecimento em base de dados sobre procedimentos assistenciais em um plano de saúde

U. R. Lamb¹, S. C. Cazella¹

¹*Programa de Pós-Graduação em Tecnologias da Informação e Gestão em Saúde, Universidade Federal de Ciências da Saúde de Porto Alegre - UFCSPA, Porto Alegre, RS*

ursulal@ufcspa.edu.br, silvioc@ufcspa.edu.br

Resumo. Introdução: O alto custo dos procedimentos, que culmina em mensalidades elevadas dos planos de saúde, a necessidade constante de fomentar programas de cuidado em saúde e a entrega ao paciente de saúde baseada em valor justificam novas abordagens em gestão, além do uso de tecnologias no gerenciamento do sistema de saúde. **Objetivo:** Aplicar o processo de descoberta de conhecimento em bases de dados (DCBD) de procedimentos de saúde buscando extrair conhecimento novo. **Método:** Aplicação do processo de descoberta de conhecimento em bases de dados com uso de aprendizado de máquina não supervisionado por meio do algoritmo *K-means* para exploração de dados. A definição de agrupamentos foi orientada pelo índice de silhueta e o teor da segregação foi analisada com base no conhecimento sobre o domínio da aplicação e dos atributos da base. **Resultado:** Foram identificados quatro agrupamentos significativos. **Conclusões:** Os agrupamentos representam conhecimento a ser explorado, podendo auxiliar no processo de tomada de decisão na gestão.

Palavras-chave: Descoberta de Conhecimento em Base de Dados. Aprendizado de máquina não supervisionado. Agrupamento de dados. *K-means*. Planos de saúde.

Abstract. Introduction: The high cost of procedures, which culminates in high monthly fees for health plans, the constant need to promote health care programs and the delivery of value-based health to the patient justify new management approaches, in addition to the use of technologies in the management of the health system. **Objectives:** Apply the process of knowledge discovery in databases (KDD) of health procedures seeking to extract new knowledge. **Method:** Application of the knowledge discovery process in databases using unsupervised machine learning algorithm *K-means* for data exploration. The definition of clusters was guided by the silhouette index and the content of segregation was analyzed based on knowledge about the application domain and the base attributes. **Result:** Four significant clusters were identified. **Conclusions:** The clustering represents knowledge to be explored, which can help in the decision-making process in management.

Keywords: Knowledge Discovery in Databases, unsupervised machine learning, clustering, K-means.

Nome do projeto: Descoberta de conhecimento em base de dados sobre procedimentos assistenciais em um plano de saúde.

1. CONTEXTO E MOTIVAÇÃO

Esta pesquisa surge de uma necessidade real para áreas de negócios que precisam analisar carteiras de beneficiários de plano de saúde. Com o aumento dos custos assistenciais, da necessidade de se melhorar a saúde das pessoas, aliadas à produção de dados pelas organizações de saúde, surge a oportunidade de se aplicarem novas soluções, como a de mineração de dados, para se extrair conhecimento acerca da base de dados de procedimentos em saúde (1). O modelo de análise obtido pode se tornar uma solução aplicável, pelos gestores, na tomada de decisão no que se refere à gestão de saúde populacional e aprimoramento do sistema de saúde (2).



1.1 Objetivo

O objetivo geral deste projeto é explorar uma base de dados sobre procedimentos em saúde provenientes da saúde suplementar por meio da aplicação do processo de DCBD, visando extrair conhecimento para apoio ao processo de tomada de decisão.

2. ATIVIDADES PRINCIPAIS

Na etapa chamada de pré-processamento, foram analisadas as bases de dados, assim como preparados os dados para aplicação do algoritmo de aprendizado de máquina não supervisionado. O algoritmo de agrupamento K-means foi utilizado e teve a função de encontrar padrões nos dados sobre procedimentos em saúde provenientes do sistema de saúde suplementar. Etapa essa, chamada de mineração de dados. Por último, foram descritos os agrupamentos obtidos e avaliados os conhecimentos explicitados, o que caracteriza a etapa de pós-processamento.

3. DESENVOLVIMENTO DO TRABALHO E RESULTADOS OBTIDOS

O projeto foi registrado sob o número 152-2022, pela Comissão de Pesquisa da Universidade Federal de Ciências da Saúde de Porto Alegre e pela sua natureza, não foi necessária a análise do Comitê de Ética. Foi aplicado o processo de DCBD visando encontrar padrões de utilização do plano de saúde pelos usuários. Esse processo inclui as etapas de seleção, limpeza, transformação, mineração dos dados e internalização. A base de dados de procedimentos em saúde provenientes do sistema de saúde suplementar, após o pré-processamento, ficou composta de 3033 instâncias e de 165 atributos. Esses procedimentos se referem à totalidade do ano de 2019. Após a aplicação do algoritmo de agrupamento K-means, foram obtidos quatro conjuntos. O agrupamento 1 foi chamado de “Pacientes com condições crônicas de saúde”, pois suas principais características são instâncias com gastos de até cinco mil reais e procedimentos relacionados à oftalmologia e à ortopedia, principalmente. O agrupamento 2, de “Pacientes com altos custos hospitalares”, pode ser associado aos usuários que tiveram gastos de internações de valores superiores a quarenta mil reais. Já o agrupamento 3, foi chamado de “Pacientes Oncológicos com internações” e o agrupamento 4, de “Pacientes Oncológicos”. O total de instância em cada grupo totalizou em 2929, 71, 15 e 18, respectivamente.

4. DESAFIOS E APRENDIZADOS

O modelo de aprendizado de máquina proposto se mostrou eficiente no agrupamento de usuários de planos de saúde de acordo com seu perfil de utilização. Novas oportunidades de descoberta de conhecimento surgem com a extensão do período da base utilizada para mais de doze meses. O principal desafio deste estudo foi conseguir reduzir a quantidade de atributos representados pelos procedimentos em saúde, pois esses totalizaram 267 atributos antes do pré-processamento, de uma maneira a não prejudicar as características de utilização da população estudada.

5. AGRADECIMENTO

A equipe do projeto agradece à Fundação de Amparo à Pesquisa do Rio Grande do Sul pelo apoio financeiro através do Projeto FAPERGS 22/2551-0000390-7 (RITE CIARS).

6. REFERÊNCIAS

1. dos Santos JRR, Dias CM, Filho AC. “Machine learning and national health data to improve evidence: Finding segmentation in individuals without private insurance”. Health Policy and Technology. (2020).
2. Yuill W, Kunz H. “Using Machine Learning to Improve Personalised Prediction: A Data-Driven Approach to Segment and Stratify Populations for Healthcare”. Studies In Health Technology And Informatics (2020).



Desenvolvimento e implementação de sistema de apoio à decisão para melhorar o controle da hipertensão arterial e do diabetes mellitus

Milena Soriano Marcolino^{1,2}, Leonardo B. Ribeiro², Breno M. Horta Melo², Caio V. Siqueira², Leandro L. Duraes², Christiane Corrêa Rodrigues Cimini³, Thiago Barbabela de Castro Soares², Antonio Luiz Ribeiro^{1,2}

¹Faculdade de Medicina, Universidade Federal de Minas Gerais (UFMG), Belo Horizonte, MG

²Centro de Telessaúde do Hospital das Clínicas da UFMG, Belo Horizonte, MG

³Universidade Federal dos Vales do Jequitinhonha e Mucuri, Teófilo Otoni, MG

milenamarc@gmail.com, leonardo.bonisson@gmail.com, brenommelo@gmail.com, caiocvs13@gmail.com, laureano_duraes@hotmail.com, christiane.cimini@gmail.com, barbabela@gmail.com, alpr1963br@gmail.com

Resumo. *Sistemas de apoio à decisão (SAD) permitem a utilização das melhores evidências científicas em favor do paciente. Em uma região brasileira com recursos limitados, a implementação de um SAD para melhorar o manejo de pessoas com hipertensão e diabetes resultou em aumento significativo da proporção de pacientes que atingiram as metas de controle para hipertensão arterial (pressão arterial sistólica < 140 mmHg e pressão arterial diastólica < 90 mmHg) e diabetes (glicohemoglobina > 8%). Os profissionais de saúde avaliaram positivamente a aplicação quanto à utilidade, usabilidade. Contudo, relataram como desafio o aumento no tempo gasto para realizar a consulta. O sistema foi aperfeiçoado e recentemente implantado em novos municípios.*

Abstract. *Decision support systems (DSS) are a way to use the best scientific evidence in favor of the patient. In a resource-constrained Brazilian region, the implementation of a DSS to improve the management of people with hypertension and diabetes resulted in a significant increase in the proportion of patients who achieved control goals for hypertension (systolic blood pressure < 140mmHg and diastolic blood pressure < 90mmHg) and diabetes (glycohemoglobin > 8%). Health professionals positively evaluated the application in terms of utility and usability. However, they reported as a challenge the increase in the time spent to carry out the consultation. The system was improved and recently implemented in other municipalities.*

Palavras-chave: Hipertensão; diabetes mellitus; sistemas de apoio a decisões clínicas.

Nome do projeto: *Control of Hypertension and Diabetes in Minas Gerais (CHARMING) (ALR).*

1. CONTEXTO E MOTIVAÇÃO

Os baixos níveis de controle da hipertensão arterial sistêmica (HAS) e do diabetes mellitus (DM) são um desafio para o Sistema Único de Saúde (SUS)(1). Este desafio exige estratégias de ação inovadoras. Sistemas de apoio à decisão (SAD) têm potencial para permitir o acesso profissional à tomada de decisões baseada em evidências mais recentes, além de fomentar a educação permanente, com potencial de melhorar o controle de HAS e DM e, com isso, reduzir as complicações e mortalidade cardiovasculares. O objetivo deste estudo foi desenvolver um SAD para o manejo de pacientes com HAS e DM, e implementá-lo em região com recursos limitados, como parte do estudo *Control of Hypertension na Diabetes in Minas Gerais (CHARMING)*.

2. DESENVOLVIMENTO DO TRABALHO

Uma equipe de médicos estabeleceu, juntamente com a equipe de tecnologia da informação do Centro de Telessaúde do Hospital das Clínicas da UFMG, os requisitos funcionais de *software* necessários para registrar e apoiar a tomada de decisões para pacientes com HAS e DM, segundo diretrizes nacionais e internacionais. A aplicação foi desenvolvida com “*design* centrado no usuário”, usando linguagem Java 1.8 e JSF 2.2, ambiente de desenvolvimento NetBeans 8.1, servidor de banco de dados PostgreSQL e *framework* Hibernate, para utilização em navegadores *Web*. A interface foi desenvolvida para ser intuitiva e autoexplicativa, utilizando Prime Faces, Bootstrap, JavaScript e jQuery. A segurança e o controle de acesso foram baseados



no Spring Security, garantindo a inviolabilidade dos dados. Houve um processo iterativo de desenvolvimento, com cooperação estreita entre equipe clínica e de desenvolvimento. Para desenvolver a funcionalidade de apoio à decisão clínica, foram revisadas recomendações de diretrizes baseadas em evidências brasileiras e internacionais. As recomendações mais relevantes foram selecionadas e foi desenvolvido um algoritmo de decisão para cada mensagem do SAD. Com essa funcionalidade, os profissionais de saúde podem acessar mensagens personalizadas para orientar o manejo de cada paciente, geradas de acordo com os dados inseridos na consulta. Cada mensagem contém recomendação resumida, texto auxiliar contendo informações complementares e referências que sustentam a recomendação. Após validação interna e externa (2), o sistema foi implementado em 10 municípios dos Vales do Jequitinhonha e Mucuri (Ataléia, Catuji, Crisólita, Frei Gaspar, Itaipé, Ladainha, Novo Oriente de Minas, Ouro Verde de Minas, Setubinha e Teófilo Otoni), como parte de intervenção multifacetada, que envolveu também educação continuada e fortalecimento das estratégias de grupo de educação em saúde. O período de acompanhamento foi de outubro/2018 a fevereiro/2021. Um questionário em escala Likert, avaliando usabilidade e satisfação, foi aplicado após 6 meses. Ao final do acompanhamento, foram realizados grupos focais com profissionais usuários.

No total, 5.078 pacientes foram acompanhados (idade mediana 56 anos [intervalo interquartil 48-62 anos], 66,4% mulheres), 27,0% apresentavam DM e 94,8% HAS. Houve aumento da proporção de indivíduos que atingiram as metas de HAS (pressão arterial sistólica < 140mmHg e pressão arterial diastólica < 90mmHg) (48,5% vs. 59,0%, $p < 0,001$) e redução significativa da proporção de pacientes com DM descontrolado (glicohemoglobina > 8%; 38,4% vs. 36,0%, $p < 0,001$). Com relação à avaliação do software, 96 profissionais responderam a enquete. Em geral, concordaram (pontuações medianas 4-5): (i) que o sistema poderia ser facilmente incorporado às rotinas de trabalho, mas os médicos alegaram que ele causava atrasos significativos nas consultas; (ii) que o sistema era de fácil compreensão e utilização; (iii) que o sistema era útil para auxiliar no manejo e melhorar o atendimento ao paciente; e ficaram globalmente satisfeitos. Os 17 participantes dos grupos focais foram consistentes ao afirmar que estavam muito satisfeitos. Contudo, o tempo gasto na consulta foi visto como um desafio para alguns participantes.

Com base nessas informações, o sistema foi aperfeiçoado, de forma a reduzir o tempo gasto e melhorar a usabilidade. Novas mensagens foram desenvolvidas e implementadas, mantendo alinhamento às diretrizes mais recentes. O sistema passou por nova validação, e foi implementado em cinco municípios dos Vales do Jequitinhonha e Mucuri (Carlos Chagas, Caraí, Poté, Novo Cruzeiro, Malacacheta) em março/2023. O acompanhamento terá a duração de 12 meses a partir da primeira consulta. Até o momento, 3930 pacientes foram cadastrados no sistema, foram realizadas 1889 consultas médicas e 839 consultas de enfermagem.

3. DESAFIOS E APRENDIZADOS

O principal desafio é a integração do sistema desenvolvido com o sistema de prontuário eletrônico usado no SUS, o que gera duplicidade de trabalho. Como perspectivas imediatas, as mensagens de suporte à decisão desenvolvidas foram revisadas e adaptadas para serem implementadas em âmbito nacional, em parceria com a Secretaria de Atenção Primária à Saúde. A perspectiva é que profissionais de saúde usuários do e-SUS APS, ao preencher dados na consulta, recebam as mensagens via notificações em campo específico.

4. REFERÊNCIAS

1. Picon RV, Dias-da-Costa JS, Fuchs FD, et al. Hypertension Management in Brazil: Usual Practice in Primary Care-A Meta-Analysis. *Int J Hypertens*. 2017;2017:1274168.
2. Marcolino MS, Oliveira JAQ, Cimini CCR, et al. Development and Implementation of a Decision Support System to Improve Control of Hypertension and Diabetes in a Resource-Constrained Area in Brazil: Mixed Methods Study. *J Med Internet Res*. 2021;23(1):e18872. doi: 10.2196/18872.
3. Viana LV, Leitão CB, Kramer CK, et al. Poor glycaemic control in Brazilian patients with type 2 diabetes attending the public healthcare system: a cross-sectional study *BMJ Open* 2013;3:e003336.



Detecção, Caracterização Inicial, Caracterização Semântica - DIS: uma nova metodologia para caracterização e análise de *drifts* temporais em dados de saúde

Bruno B. Paiva¹, Claudio Andrade¹, Gabriel Moraes¹, Ricardo Cardoso³, Milena S. Marcolino¹, Ana Etges³, Polianna Delfino-Pereira¹, Carisi Polanczyk², Marcos A. Gonçalves¹, Leonardo Rocha²

¹Universidade Federal de Minas Gerais, Belo Horizonte, MG, ²Universidade Federal de São João del-Rei, São João del-Rei, MG, ³Universidade Federal do Rio Grande do Sul, Porto Alegre, RS

Resumo: Neste trabalho, propomos uma metodologia denominada DIS (detecção, caracterização inicial, caracterização semântica), que por meio de técnicas inspiradas em recentes estudos sobre processamento de linguagem natural, é capaz de analisar mudanças em desfechos e variáveis de saúde ao longo do tempo. Avaliamos a metodologia em dois conjuntos de dados (Registro Hospitalar Multicêntrico Nacional de Pacientes com Covid-19 e MIMIC-IV), clarificando o porquê da queda na mortalidade em ambos os conjuntos, destacando o papel da vacinação na pandemia da covid-19 e a diminuição de eventos iatrogênicos no conjunto de dados MIMIC-IV. A compreensão destes eventos podem contribuir para melhorar a eficácia dos algoritmos de aprendizado de máquina, aplicados aos dados de assistência médica, além de ajudar a entender quais fatores influenciam os desfechos dos pacientes ao longo do tempo e o impacto das novas tecnologias e políticas nos mesmos.

Abstract : In this work, we propose a methodology called DIS (detection, initial characterization, semantic characterization), which, using techniques inspired by recent studies on natural language processing, is capable of analyzing changes in outcomes and health variables over time. We evaluated the methodology in two sets of data (National Multicentric Hospital Registry of Patients with Covid-19 and MIMIC-IV), clarifying the reason for the drop in mortality in both sets, highlighting the role of vaccination in the covid-19 pandemic and the decrease in iatrogenic events in the MIMIC-IV dataset. Understanding these events can contribute to improving the effectiveness of machine learning algorithms applied to healthcare data, as well as helping to understand which factors influence patient outcomes over time and the impact of new technologies and policies on them.

Palavras-chave: Desfechos em saúde; *drifts* temporais; processamento de linguagem natural.

Nome do projeto: Predição de Desfechos Clínicos e Econômicos por Meio de Representações Semânticas Multi-modais de Pacientes Resilientes a *Drifts* Temporais (Marcos André Gonçalves)

1. CONTEXTO E MOTIVAÇÃO

Dados de saúde são um recurso valioso na busca por intervenções capazes de melhorar os desfechos de pacientes, assim como para entender mudanças de padrão que porventura tenham acontecido. Trabalhos como (Paiva, 2022)¹ identificaram mudanças de padrão relevantes na mortalidade hospitalar de pacientes com covid-19, provenientes de coorte brasileira multicêntrica, sugestivos do impacto da vacinação e da mudança no perfil da doença com o aparecimento de novas variantes. *Drifts* desse tipo também ocorreram com outras doenças, como câncer de mama, que passou a ser um tumor praticamente curável na maioria dos casos, com o advento da hormonioterapia e dos imunobiológicos. No panorama atual, a relevância e o potencial de detectar este tipo de mudança cresceu, devido aos avanços na capacidade de processamento e armazenamento de dados. Neste contexto, propomos um arcabouço metodológico inspirado em processamento de linguagem natural (PLN), para caracterizar, analisar e compreender os principais determinantes de um dado desfecho clínico e sua evolução ao longo do tempo.

2. ATIVIDADES PRINCIPAIS

O trabalho foi realizado por alunos de pós-graduação e iniciação científica com participação dos orientadores.



3. DESENVOLVIMENTO DO TRABALHO

Nossa proposta é dividida em três passos. O primeiro é a **detecção**, que decide “se” devemos prosseguir com a análise, avaliando “se” os dados têm variações temporais relevantes (ou não). Esse passo é conduzido ao avaliar o desvio de distribuição das variáveis ao longo do tempo, usando métricas como a divergência de Shannon-Jensen, ou ao observar desvios significativos de performance em métricas de efetividade de classificadores (como acurácia, *score* F1 e *recall*). O segundo é a **caracterização inicial**, que visa obter uma compreensão geral do “o que(s)” e “como (s)” em relação às mudanças nos dados analisados. Nele, buscamos por mudanças de padrão em geral e na distribuição do desfecho de interesse (ex. houve um aumento/diminuição de óbitos ou de necessidade de ventilação mecânica nos pacientes?) que possam guiar as análises no passo seguinte, denominado **caracterização semântica**. Nesse terceiro e último passo, estamos preocupados em aprender “por que” as mudanças que descobrimos no passo anterior ocorreram. Para este fim, usamos as técnicas de PNL como a geração de *embeddings* e comparação de distâncias semânticas entre vetores de desfechos clínicos, para obter vetores capazes de representar semanticamente nossos desfechos ao longo do tempo e então caracterizá-los e compará-los. Atualmente estamos utilizando duas bases de dados, do projeto “Registro hospitalar multicêntrico nacional de pacientes com covid-19”² (cerca de 10 mil pacientes) e da MIMIC-IV (cerca de 60 mil pacientes). O estudo foi aprovado pela Comissão Nacional de Ética em Pesquisa (CAAE 30350820.5.0000.0008).

4. DESAFIOS E APRENDIZADOS

A Figura 1-A mostra um exemplo de grafo construído para dois pacientes. Este grafo é a base para a construção de um *embedding* semântico das entidades, incluindo os desfechos. A Figura 1-B mostra as principais mudanças de padrão entre 2020 e 2021, obtidas ao analisar as diferenças entre os *embeddings* dos desfechos. O desfecho “óbito” perdeu similaridade de 2020 para 2021 com as faixas etárias mais altas e ganhou similaridade com outros sinais clínicos e laboratoriais, como FiO_2 mais baixo à admissão. Esse *drift* provavelmente está relacionado com a vacinação no Brasil, que começou com faixas etárias mais altas.



Figura 1. Resultados parciais da análise semântica na base da covid-19.

O principal desafio do projeto é a obtenção de outras bases para análise, devido aos requisitos e cuidados de privacidade na saúde. Isso leva à dificuldade de encontrar dados adequados, como texto livre de prontuários ou exames de imagem. Além disso, a segunda dificuldade reside no desenvolvimento de uma metodologia genérica que possa ser aplicada a dados de diferentes naturezas, como a diferença de granularidade temporal entre a coorte multicêntrica de pacientes com covid-19 e a base MIMIC-IV.

5. REFERÊNCIAS

1. Paiva, B, Pereira, P, Gomes, V, Silva, M, Valiense, C, Marcolino, M, Gonçalves, Characterizing and Understanding Temporal Effects in COVID-19 Data. In Workshop on Healthcare AI and COVID-19, 2022.



Inteligência Artificial para a Saúde no TikTok: Uma avaliação da informação

Giovanna M. Vilas Boas¹, Teresa Vitória C. Rocha¹, Isabella L. Ferreira¹, Mylena Maria G. de Almeida¹, José Artur C. Fernandes¹, Ana Paula C. da Silva¹, Cristiane dos Santos Dias¹, Zilma Silveira Nogueiras Reis¹

¹Universidade Federal de Minas Gerais, Belo Horizonte, MG

{giovannamvb, teresavitoriacr, isabellalellism, mylenamga2018, jose.artur.crquer.fernandes}@gmail.com, ana.coutosilva@dcc.ufmg.br, {profacristianedias, zilma.medicina}@gmail.com

Resumo. O TikTok é uma rede social de crescimento rápido que tem sido usada para divulgar conteúdos diversos, incluindo tópicos sobre Inteligência Artificial (IA) na área da Saúde. Para avaliar como o tema da IA na Saúde está sendo apresentado no TikTok, 93 vídeos foram avaliados manualmente por voluntários. Este estudo inicial mostrou o potencial do TikTok como aliado da divulgação tanto de informação de qualidade, como também de desinformação, corroborando a importância de estratégias para Divulgação Científica de qualidade.

Abstract. TikTok is a fast-growing social network that has been used to disseminate various content, including topics related to Artificial Intelligence (AI) in the field of Healthcare. With the aim of evaluating how the topic of AI in Health is being presented on TikTok, 93 videos were manually evaluated by three different assessors. This preliminary study demonstrated the potential of TikTok as an ally in promoting both high-quality information and misinformation, underscoring the importance of strategies for quality Scientific Communication.

Palavras-chave: Inteligência Artificial, Saúde, Redes Sociais, Divulgação do Conhecimento

1. CONTEXTO E MOTIVAÇÃO

Atualmente, as redes sociais são utilizadas como fonte de informações em diferentes temas, entre eles conteúdos relacionados à Saúde. Entre as diversas redes sociais, o TikTok, plataforma de vídeos curtos, tem atraído um grande número de participantes e contava com 1,051 bilhões de usuários em janeiro de 2023 (1). Entre o grande volume de vídeos compartilhados, uma quantidade significativa deles aborda o uso da Inteligência Artificial na Saúde. No entanto, ao mesmo tempo em que essa ferramenta pode ser uma aliada na alfabetização do público leigo em relação a Inteligência Artificial na Saúde, o grande fluxo de vídeos criados sem uma moderação quanto a qualidade da informação divulgada também traz a preocupação quanto a desinformação vinculada a esses materiais (2). Assim, trabalhos que analisam o tipo de informação que é amplamente disseminada nesta rede são extremamente relevantes para, por exemplo, modelos de moderação mais assertivos que evitem a disseminação de conteúdo de desinformação. Este trabalho tem como objetivo apresentar um estudo inicial dos vídeos mais populares que abordam o tema do uso da Inteligência Artificial na Saúde. Para esta análise, 93 vídeos foram avaliados manualmente por um conjunto de voluntários, a partir de um questionário que examinou a autoria, natureza, acessibilidade e confiabilidade do conteúdo.

2. METODOLOGIA E RESULTADOS

O TikTok retorna em suas pesquisas os vídeos mais relevantes, classificados de acordo com algoritmo próprio da rede social. Os critérios para a categorização são baseados em características do vídeo e dados de entrada do usuário. Seguindo estudo anterior (2), foram coletados os 100 primeiros vídeos mais relevantes retornados considerando as palavras-chaves Inteligência Artificial na Saúde, Saúde Digital e Inteligência Artificial na Medicina, totalizando 300 vídeos. Para cada vídeo, foram coletados o número de visualizações, curtidas, comentários e hashtags. Após a coleta, foi realizada uma inspeção manual por 3 voluntários que excluíram das análises os vídeos em outros idiomas, duplicados, indisponíveis, ou que fugiam do tema de interesse. Dos 300 vídeos coletados, apenas 93 se encaixaram no critério de inclusão e foram avaliados. A

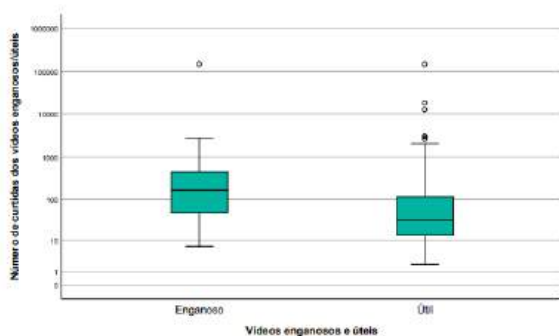


Figura 1 apresenta a nuvem de palavras com as 50 hashtags mais frequentes associadas aos vídeos analisados.

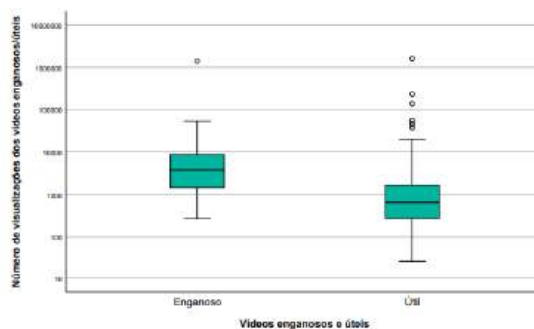


Figura 1 - Nuvem de palavras das Hashtags dos vídeos avaliados.

Com base na avaliação manual dos avaliadores, os vídeos foram classificados em três grupos distintos: úteis (n=68), enganosos (n=14) e neutros (n=11). O número de visualizações e curtidas para os vídeos classificados nas categorias útil e enganoso foram representados por boxplot (Figura 2), com o intuito de analisar as características de popularidade desses grupos. A distribuição apresentada mostra que vídeos que possuem informações enganosas apresentam um alcance maior, corroborando trabalhos anteriores que mostraram que conteúdo de desinformação é difundido mais rapidamente e possui maior alcance (3).



(a)



(b)

Figura 2 : Número de curtidas (a) e visualizações (b) para vídeos enganosos e úteis.

4. DESAFIOS E APRENDIZADOS

Os resultados aqui apresentados corroboram a crescente preocupação quanto à desinformação e a importância da divulgação científica de qualidade. A promoção de conteúdos de qualidade, que engajem o público e que tenham uma linguagem acessível é um dos grandes desafios da Era Digital.

5. REFERÊNCIAS

1. We are Social, Meltwater. Digital 2023: Global Overview Report. 2023. <https://www.meltwater.com/en/global-digital-trends>
2. Yeung A, Ng E, Abi-Jaoude E. TikTok and Attention-Deficit/Hyperactivity Disorder: A Cross-Sectional Study of Social Media Content Quality. The Canadian Journal of Psychiatry. 2022.
3. Vosoughi S, Roy D, Aral S. The spread of true and false news online. Science. 2018 Mar 9;359(6380):1146–51.



Anais do I Simpósio CI-IA Saúde da UFMG

Passado, presente e futuro.

25 de setembro de 2023 - Belo Horizonte - MG



microRNAs in cerebrospinal fluid associated with alzheimer's disease: a systematic review and pathway analysis using a machine learning approach

Jessica Diniz Pereira ^{1,2} Lívia Cristina Ribeiro Teixeira ¹, Izabela Mamede Costa Andrade da Conceição ³, Paulo Caramelli ², Marcelo Rizzatti Luizon ³, Adriano Alonso Veloso ⁴, Karina Braga Gomes ^{1,2}

¹Faculdade de Farmácia, Universidade Federal de Minas Gerais, Belo Horizonte, MG

²Faculdade de Medicina, Universidade Federal de Minas Gerais, Belo Horizonte, MG

³Instituto de Ciências Biológicas, Universidade Federal de Minas Gerais, Belo Horizonte, MG

⁴Instituto de Ciências Exatas, Universidade Federal de Minas Gerais, Belo Horizonte, MG

jessicadinizpereira@hotmail.com, liviacrteixeira@gmail.com, iza.mamede@gmail.com,
paulo.caramelli@gmail.com, luizonmr@gmail.com, adrianov@dcc.ufmg.br, karinabgb@gmail.com

Abstract: The increase in life expectancy and population aging are related to the increase in Alzheimer's disease (AD) frequency, since aging is a risk factor for the development of the disease. Currently, the diagnosis of AD is based on clinical evaluation. However, due to the complexity of the condition and the need for more precise methods, the search for biomarkers becomes necessary. MicroRNAs (microRNAs), small strings of non-coding RNA that regulate cellular processes, are the subject of growing interest in relation to AD, since several miRNAs have been identified as markers for the diagnosis and progression of this disease. The present study aimed to perform a search for miRNAs differentially expressed in cerebrospinal fluid (CSF) of patients with AD compared to cognitively compatible individuals, from a public database and using machine learning tools, with validation of results in a systematic review and proposition of regulated biological pathways. The results suggest that miRNA-1274a, miRNA-193a-5p, miRNA-28-3p, miRNA-30a-3p, miRNA-145, miRNA-19b and miRNA-143 may be involved in biological pathways related to the disease and may be therapeutic targets for AD in the future.

Keywords: Alzheimer's disease, miRNAs, biomarkers, systematic review, biological pathways, machine learning.

1. BACKGROUND AND MOTIVATION

Approximately 55 million people have dementia, and 10 million new cases are diagnosed worldwide each year (1). Alzheimer's disease (AD) is the most common type of dementia, accounting for 60% to 70% of reported cases (1). AD has been identified as one of the 10 most disabling and lethal diseases globally (2). Currently, there is no cure for AD, and its pathophysiology has not yet been fully elucidated (3). For a long time, the diagnosis of AD was exclusively based on clinical criteria, with neuropathological evaluation reserved solely to confirm the diagnosis (4). MicroRNAs (miRNAs) are small non-coding RNAs composed of 21 to 25 nucleotides, playing a significant role in gene modulation and several molecular mechanisms (5). Several studies have associated the miRNAs to AD pathophysiology (5). miRNAs can be found in plasma, tissues, and cerebrospinal fluid (CSF), making them targets of study as potential AD biomarkers (5).

1.1 Objective

This study aimed to perform an in silico analysis of miRNA expression in the CSF of patients with AD compared to cognitively healthy controls and to perform an analysis of the biological pathways regulated by these miRNAs.



2. ACTIVITIES

Perform a search for public datasets in the Gene Expression Omnibus (GEO) that assess miRNAs in the CSF of patients with AD and controls. Subsequently, apply a machine learning (ML) algorithm to identify differentially expressed miRNAs that may be associated with AD. Conduct a systematic literature review (SLR) to characterize the miRNAs in CSF involved in AD. Perform an integrative analysis using the SLR data to validate the miRNAs identified in the ML research. Next, conduct an enrichment analysis and characterization of pathways related to the identified miRNAs, focusing on their biological functions and known associations with AD.

3. WORK DEVELOPMENT

Using the LightGBM machine learning algorithm, we selected miRNAs from GEODATA with the best predictive values and compared them with miRNAs identified in the systematic review. The systematic review initially identified 2060 studies, from which 24 met the inclusion criteria after removing duplicates and reviewing titles, abstracts, and full texts. Seven miRNAs were common in both approaches: 1274a, 193a-5p, 28-3p, 30a-3p, 145, 19b, and 143, all showing reduced expression in AD patients compared to controls. These miRNAs are involved in key biological pathways, including TGF-beta signaling, MAPK, ERBB, and EGF pathways and ATM-dependent DNA damage pathway, all linked to the pathophysiology of AD.

4. CHALLENGES AND LEARNINGS

The novel findings presented in this study hold promise in aiding the development of early diagnostic tests for AD. Moreover, the exploration of these signaling pathways can contribute to a deeper understanding of the pathophysiology of AD, which can potentially uncover new insights and therapeutic targets to AD.

5. REFERENCES

1. World Health Organization (WHO). Dementia [Internet]. Geneva: WHO; 2022 [updated 2022 Jan 17; cited 2023 May 4]. Available from: <https://www.who.int/news-room/fact-sheets/detail/dementia>
2. Pan American Health Organization (PAHO). WHO reveals leading causes of death and disability worldwide between 2000 and 2019 [Internet]. Washington, D.C.: PAHO; 2020 Dec 9 [cited 2023 May 4]. Available from: <https://www.paho.org/en/news/9-12-2020-who-reveals-leading-causes-death-and-disability-worldwide-between-2000-and-2019>
3. Alzheimer's Disease International. Dementia facts and figures [Internet]. Available from: <https://www.alzint.org/about/dementia-facts-figures/>. Accessed on: May 5, 2023.
4. Trejo-Lopez JA, Yachnis AT, Prokop S. Neuropathology of Alzheimer's disease. Neurotherapeutics. 2021;18:1-13. doi: 10.2174/1570159X18666200528142429. PMID: 32484110; PMCID: PMC7709159.
5. Lukiw WJ. Micro-RNA speciation in fetal, adult and Alzheimer's disease hippocampus. Neuroreport. 2007;18(3):297-300.



Perfil de leitos brasileiros em tempos pré, durante e pós pandêmicos

Lucas V. Cortez¹, Alessandro Bigoni, Ramon¹ Pereira(orientador)¹

¹IQVIA, São Paulo SP

lucas.cortez@iqvia.com, alessandro.bigoni@iqvia.com, ramon.pereira@iqvia.com

Resumo. Entender o SUS e seus sistemas de informações ampliam nosso escopo de tomada de decisões gerenciais relacionados a saúde pública. Cada estabelecimento de saúde é visto como uma unidade informacional desses sistemas e são importantes para o entendimento sobre estrutura, organização e operacionalização de todo o sistema de saúde assistencial. O período pandêmico ressaltou a importância da análise crítica e séria de dados relacionados a saúde, principalmente pelos altos número de ocupações hospitalares que extrapolaram os limites estruturais do SUS. Dessa maneira objetivamos avaliar o perfil de leitos pré, durante e pós pandemia para todos os estados brasileiros. Após a avaliação de leitos percebemos que os números de 2019 não foram suficientes para o suporte eficiente à saúde pública ao grande número de ocupação causado no período pandêmico e a maior parte dos nossos números hoje se aproximam daqueles vistos em 2019, distanciando do ótimo definido pela OMS.

Palavras-chave: Leitos hospitalares; SUS, manipulação de dados; pandemia.

1. INTRODUÇÃO

Os estabelecimentos de saúde (ES) brasileiros são definidos pelo Ministério da Saúde (MS) como “espaços físicos delimitados e permanentes onde são realizadas ações e serviços de saúde humana sob responsabilidade técnica”, definição que elucida o que é a unidade estrutural geográfica do Sistema Único de Saúde (SUS). Esses estabelecimentos possuem naturezas jurídicas, recursos físicos e humanos distintos, assim como suas características de gerenciamento são registradas em seus cadastros nacionais de estabelecimento de saúde (CNES). O CNES além de cadastro de um ES no SUS, é também base para a maior parte dos sistemas informacionais de saúde pública, por disponibilizar informações de infraestrutura e operacionalização [1,2].

A coleta e o processamento de dados de saúde são importantes para tomadas de decisões e após a pandemia de COVID19 mostraram-se ainda mais importantes. O aumento ocupacional sentido em cada ES [3] combinado com a distribuição desigual da capacidade de leitos [4] montou um cenário crítico no Brasil, ao SUS e sua estrutura. Somado a pandemia, a falta de planejamento [4], resultou em uma catástrofe na saúde pública brasileira com ocupações muito mais altas que a disponibilidade de leitos para a população. No entanto, os sistemas informacionais do SUS levantaram grande número de dados durante o evento pandêmico, auxiliando pesquisadores e legisladores em tomadas de decisões mais assertivas.

Desse modo, o objetivo deste trabalho é descrever o cenário do sistema assistencial de saúde no Brasil desde 2019 até 2023 e compará-lo com o cenário de ocupação de leitos durante a pandemia.

2. MÉTODOS

Selecionamos oito bases de dados pertencentes ao eSUS Notifica–Módulo de internações, cinco referentes às séries históricas de 2019 (ano anterior à pandemia) até 2023 (ano posterior à pandemia) de disponibilização de leitos hospitalares por CNES e outras três bases referentes aos registros de ocupações e saídas(óbitos/altas) hospitalares nos anos de pandemia (2020 até 2022). Após a seleção do banco de dados, partimos para a etapa de triagem de cada um dos conjuntos de dados selecionados: leitos e ocupações. Sobre os leitos, consideramos todos os leitos disponíveis como a capacidade máxima de suporte do estabelecimento de saúde, e para ocupação consideramos a diferença diária entre todos os tipos de entradas e saídas hospitalares. Pouca ou nenhuma inconsistência foi encontrada, o que não aconteceu para os dados de ocupações. Essas inconsistências foram ponderadas a partir do CNES com a capacidade máxima no estado, logo se a ocupação hospitalar em um CNES supera a maior capacidade estrutural presente no estado esse outlier era pontuado e substituído pela média móvel da ocupação naquele estabelecimento com uma



amplitude de 10 dias corridos. Método utilizado também para as saídas hospitalares, tanto óbitos quanto altas.

A terceira etapa consistiu na manipulação e análise descritiva dos dados. Essa avaliação foi feita a nível de CNES, estado e região, enfatizando o nível estadual. Aqui, reportamos a capacidade de suporte de um estado através da média anual da quantidade de leitos no estado a cada 1000 habitantes a nível anual. Média essa que foi ponderada por mês, uma vez que meses terminais apresentam menor importância na capacidade de suporte anual. Para as ocupações calculamos a diferença entre as saídas e as entradas num determinado CNES e somamos mensalmente, obtendo a ocupação mensal naquele estabelecimento. Desse modo, conseguimos comparar como as ocupações hospitalares afetam a capacidade de suporte a nível estadual.

3. RESULTADOS E DISCUSSÃO

Após a triagem dos dados, obtivemos dois conjuntos de dados referentes aos leitos e ocupações hospitalares: com aproximadamente 378.312 linhas e 1.591.611 linhas respectivamente. Vale ressaltar que para a base de ocupação não houve notificações para o estado de Espírito Santo, inviabilizando as análises para esse estado.

A despeito da disponibilidade de leitos hospitalares e capacidade de suporte brasileira, percebemos um aumento no número de leitos devido a pandemia em comparação com o período pré-pandêmico, no entanto, logo após o término o número de leitos tende a aproximar-se dos números pré-pandêmicos. Quando observamos a nível estadual percebemos que os estados da região Norte, principalmente Amazonas (1.63 lto/hab), Amapá (1.71 lto/hab) e Pará (1.83 lto/hab) são aqueles com situações mais críticas quanto a disponibilidade de leitos. e no outro extremo temos o Distrito Federal (4.65 lto/hab) e Santa Catarina (4.53) com maiores números de disponibilidades de leitos por habitantes em 2023.

Ao analisarmos a ocupação de leitos nos anos de pandemia, percebemos o tamanho do estresse para o sistema de saúde, chegando a ocupações 8 vezes maiores que a suportada estruturalmente pelo estado, como visto em Tocantins, e em média 2 vezes maior que a suportada pela estrutura do SUS. Dito isso, a redução na proporção de leitos vista para o ano de 2023 não seria esperada, visto que é sabido que durante uma catástrofe pandêmica havia a necessidade média de 2 vezes mais leitos em cada estado.

4. CONCLUSÃO

Hoje a capacidade de suporte de leitos no Brasil está em declínio e mostrou-se maior que aquela vista no período pré pandêmico, porém longe dos números preconizados pela OMS. O que enfatiza a necessidade de uma melhor gestão dos números de leitos no território brasileiro.

5. REFERÊNCIAS

- [1] Ministério da Saúde. Matrizes de Consolidação - Estabelecimento de saúde. Acessado em: 23/08/2023, Disponível em: bvsms.saude.gov.br/bvs/as_udelegis/gm/2017/MatrizesConsolidacao_comum/251439.html#:~:text=Estabelecimento de Saúde
- [2] Ministério da Saúde. Cadastro Nacional de Estabelecimentos de Saúde - CNES. pub.2018. Acessado em: 23/08/2023. Disponível em: <https://cnes-cadastro-nacional-de-estabelecimentos-de-saude-ministerio-da-saude>
- [3] Costa Lino, D. Barreto, R. Souza, D. et al. Impact of lockdown on bed occupancy rate in a referral hospital during the COVID-19 pandemic in northeast Brazil. *Braz J Infect Dis*, 24-5, 2020.
- [4] Matarazzo, H. Zoca, B. et al. Cenário dos Hospitais no Brasil - 2019. CN Saúde, 2019.



Predição de comportamento suicida em crianças e adolescentes em uma instituição de cuidados psiquiátricos

Isis F. Carvalho¹, Debora Miranda², Ana Paula Couto da Silva¹, Anísio M. Lacerda¹, Wagner Meira Jr.¹,
Marco A. Romano-Silva², Maria Carolina Lobato², Gisele L. Pappa¹

¹Departamento de Ciência da Computação, Universidade Federal de Minas Gerais, Belo Horizonte, MG

²Faculdade de Medicina, Universidade Federal de Minas Gerais, Belo Horizonte, MG

{isisferreira, ana.coutosilva, anisio, meira, glpappa}@dcc.ufmg.br

{debora.m.miranda, romanosilva, mcarollobato}@gmail.com

Resumo. O suicídio já é a quarta principal causa de morte nos indivíduos na faixa entre 15-19 anos, de acordo com a OMS. Preditores de comportamentos relacionados ao suicídio podem ajudar na tarefa de intervir para evitar ou monitorar riscos futuros de suicídio. Neste artigo, foi analisada uma amostra de indivíduos atendidos num centro psiquiátrico infantil, localizado em Belo Horizonte. Algoritmos de aprendizado de máquina foram usados para gerar modelos para prever o comportamento suicida, e as características que melhor explicam esse comportamento complexo também foram analisadas. Os resultados preliminares mostram que o modelo de Random Forest é o que, como observado em trabalhos relacionados, apresentou melhor performance quanto às métricas de sensibilidade e especificidade, identificando com alta precisão tanto casos positivos quanto negativos dos três desfechos analisados neste estudo: automutilação, ideação suicida e tentativa de suicídio.

Abstract. Suicide is already the fourth leading cause of death for individuals aged 15-19, according to the WHO. Predictors of suicide-related behaviors can help in the task of intervening to prevent or monitor future suicide risks. In this article, a sample of individuals treated at a child psychiatric center located in Belo Horizonte was analyzed. Machine learning algorithms were used to generate models to predict suicidal behavior, and the features that best explain this complex behavior were also analyzed. The preliminary results show that the Random Forest model is the one that, as observed in related works, presented the best performance in terms of sensitivity and specificity metrics, accurately identifying both positive and negative cases of the three outcomes analyzed in this work: self-harm, suicidal ideation, and suicide attempt.

Palavras-chave: Machine Learning; Predição de Suicídio; Fatores de Risco Suicida; Psiquiatria da Criança e Adolescente.

Nome do projeto: Estudo de associação e predição de desfechos psiquiátricos graves em crianças e adolescentes por meio de inteligência artificial (orientadora Gisele L. Pappa).

1. CONTEXTO E MOTIVAÇÃO

A cada 40 segundos, uma pessoa morre por suicídio no mundo, [1], sendo o suicídio a maior causa de morte de crianças e adolescentes em países desenvolvidos [1]. Dados os impactos da morte de crianças e adolescentes na estrutura das famílias e da própria sociedade, a investigação dos fatores que levam um indivíduo a apresentar comportamentos suicidas é de extrema relevância. Assim, o objetivo deste trabalho é identificar variáveis preditoras do comportamento suicida, de forma a auxiliar na proposta de intervenções direcionadas em casos de pacientes que apresentam alto risco de suicídio

Grande parte dos trabalhos da área se baseiam em dados provenientes de estudos populacionais longitudinais. O trabalho aqui apresentado baseia-se em [2], no qual os autores utilizam modelos de aprendizado de máquina para predição de risco suicida em jovens adultos da Escócia. Diferentemente de [2], nossas análises de predição de risco suicida em uma população em situações de crises de saúde mental utilizam dados originados a partir da digitalização de registros físicos coletados durante os atendimentos de emergência no Centro Psíquico de Adolescência e de Infância (CEPAI), associado à rede de hospitais da Fundação



Hospitalar do Estado de Minas Gerais (FHEMIG). Dito isso, neste trabalho, serão propostos modelos para predição de três variáveis alvo distintas: automutilação, ideação suicida e tentativa de suicídio.

2. ATIVIDADES PRINCIPAIS

A metodologia utilizada para a proposta do modelo preditivo segue as etapas principais da solução de problemas que aplicam modelos de aprendizado. A primeira etapa consiste na exploração e limpeza dos dados. Foram analisados 2.365 registros de pacientes e 150 atributos (variáveis) associados a cada um deles. Desses 150 atributos, 57 foram escolhidos com a ajuda de especialistas para serem usados nos modelos, passando por uma etapa de one hot encoding. Ao final desse processo, e após a remoção de 74 registros não informativos (também com a ajuda de especialistas), o conjunto de dados se reduziu a 2.291 registros com 154 variáveis. Em seguida, modelos tradicionais de aprendizado de máquina (Regressão Logística, Random Forest e Gradient Boosting) foram treinados e testados a partir dos dados pré-processados e seus desempenhos foram comparados através de diferentes métricas: área sob a curva ROC (AUC ROC), valor preditivo positivo (PPV), sensibilidade e especificidade. A métrica de avaliação mais valorizada neste trabalho foi a sensibilidade, que é a quantidade de casos positivos que o modelo identifica corretamente, entre todos os casos positivos. Escolhemos essa métrica pois, no contexto de prevenção ao suicídio, é preferível que tenhamos mais falsos positivos do que falsos negativos (ou seja, é preferível prestar mais atenção em pacientes de menor risco do que negligenciar pacientes que apresentam alto risco de comportamento suicida).

Várias configurações experimentais foram testadas, de forma a comparar as diferentes performances, incluindo o uso de oversampling no conjunto de treinamento, a seleção de subconjuntos de features para alimentar o modelo, e o ajuste de parâmetros específicos de cada modelo.

3. DESENVOLVIMENTO DO TRABALHO

O modelo de melhor desempenho, considerando as métricas de sensibilidade e especificidade foi o Random Forest com oversampling em cima de um subconjunto das variáveis da base (escolhidas com o auxílio de especialistas). Esse subconjunto é composto basicamente pelos motivos pelos quais ajuda médica no CEPAI foi procurada, pelos diagnósticos recebidos pós-atendimento, e por algumas informações pessoais, como idade, sexo, se o paciente teve que ser hospitalizado, se o paciente é recorrente no CEPAI, e com quantas pessoas o paciente mora. Essa configuração experimental foi a que maximizou o poder preditivo para todos os desfechos testados.

Dentre os 5 fatores mais importantes para as três tarefas, depressão e sexo sempre se mostraram presentes, o que nos diz que ambos os fatores demandam atenção especial. O fato da depressão ter sido um dos preditores mais importantes em vários experimentos feitos durante o projeto reforça que esse fator é chave para o tratamento dos pacientes pelos especialistas.

4. DESAFIOS E APRENDIZADOS

Os dados analisados, por terem sido coletados manualmente, possuem diversos erros de preenchimento, o que dificultou a etapa de tratamento de dados. Ademais, a quantidade de registros (e a pequena proporção de casos positivos de suicídio em relação ao total) limita as classes de modelos a serem aplicadas para responder às tarefas propostas. Em trabalhos futuros, pretende-se obter dados inseridos em outros contextos, para enriquecer nossos dados e validar os modelos propostos em outras distribuições. Pretende-se também utilizar causalidade para entender as intervenções possíveis de serem feitas para evitar o trágico desfecho de suicídio.

5. REFERÊNCIAS

1. Knipe, Padmanathan et. al. "Suicide and self-harm". Lancet 399 (2022).
2. Van Mens, de Schepper et. al. "Predicting future suicidal behaviour in young adults, with different machine learning techniques: A population-based longitudinal study". Journal of Affective Disorders 271 (2020).



Predição de Necessidade de Internação em Terapia Intensiva e Custos de Tratamento de Pacientes com COVID-19

Claudio M.V. Andrade¹, Bruno Paiva¹, Gabriel Morais¹, Ricardo Cardoso³, Milena Marcolino¹, Ana Etges³, Polianna Delfino-Pereira¹, Carisi Polanczyk³, Marcos A. Gonçalves¹, Jussara Almeida¹, Leonardo Rocha²

¹Universidade Federal de Minas Gerais, Belo Horizonte, MG, ²Universidade Federal de São João del-Rei, São João del-Rei, MG, ³Universidade Federal do Rio Grande do Sul, RS

claudio.valiense@dcc.ufmg.br

Resumo: Este estudo tem por objetivo a predição de necessidade de internação em unidade de terapia intensiva (UTI) e custos hospitalares de pacientes com covid-19. Para tanto, ele explora dados de pacientes de cinco hospitais, entre março e agosto de 2020. Analisamos 12 características demográficas, clínicas e custo total. Observou-se variações significativas no custo médio por paciente entre os hospitais, sendo o uso da UTI o principal fator de impacto no custo final. Foram investigadas diferentes tarefas de predição, bem como diferentes estratégias de pré-processamento para reduzir o impacto da grande variabilidade e pequeno tamanho do conjunto de dados. Os resultados moderados de predição obtidos apresentaram diferenças significativas de acurácia da predição de UTI nos diferentes hospitais, bem como na predição da faixa de custo do paciente. Adicionando-se informações dinâmicas relacionadas ao tratamento recebido na UTI, observamos melhorias significativas na acurácia da predição da faixa de custo, em todos hospitais.

Abstract: This study aims to predict the need for Intensive Care Unit (ICU) admission and hospital costs for patients with COVID-19. To achieve this, it explores patient data from five hospitals between March and August 2020. We analyzed 12 demographic, clinical, and total cost features. Significant variations in the average cost per patient were observed among the hospitals, with ICU usage being the primary cost impact factor. Various prediction tasks were investigated, along with different preprocessing strategies to mitigate the impact of the dataset's substantial variability and small size. The achieved moderate prediction results showed significant differences in ICU prediction accuracy among the different hospitals, as well as in predicting the patient's cost range. By incorporating dynamic information related to the treatment received in the ICU, we observed substantial improvements in the accuracy of predicting the cost range across all hospitals.

Palavras-chave: modelos de predição de custo; covid-19; unidade de terapia intensiva; cuidados críticos.

Nome do projeto: Predição de Desfechos Clínicos e Econômicos por Meio de Representações Semânticas MultiModais de Pacientes Resilientes a Drifts Temporais (Marcos André Gonçalves).

1. CONTEXTO E MOTIVAÇÃO

Esse estudo está contextualizado na primeira onda da pandemia de covid-19, quando os sistemas de saúde foram sobrecarregados e enfrentaram desafios significativos para fornecer atendimento médico adequado aos pacientes[1]. Além, obviamente, da importância da vida das pessoas, outro fator importante está relacionado ao impacto da pandemia na viabilidade econômica dos hospitais, ou seja, como se deu a relação entre os diferentes tratamentos proporcionados aos pacientes e seus custos.

1.1 Objetivo

Investigamos duas tarefas de predição: primeiramente uma de classificação binária para predizer, tomando como entrada apenas informações disponíveis no momento de admissão do paciente, se ele iria para a UTI. Esta tarefa foi motivada por um resultado da análise de dados que apontou ser este o atributo mais importante para a definição dos custos totais. A segunda tarefa de classificação multi-classe objetivou prever a faixa dos custos totais (custo baixo, moderado ou alto) do paciente.

2. ATIVIDADES PRINCIPAIS

O trabalho foi realizado por alunos de pós-graduação e iniciação científica com participação dos orientadores. O estudo foi aprovado pela CAAE 30350820.5.0000.0008.



3. DESENVOLVIMENTO DO TRABALHO

Inicialmente, realizou-se uma análise detalhada de dados referentes a pacientes com covid-19, entre os meses de março a agosto de 2020, buscando identificar padrões, tendências e fatores que influenciaram o custo hospitalar total do paciente. A amostra incluiu pacientes admitidos em quatro hospitais públicos/acadêmicos e um privado, localizados no Rio Grande do Sul e em Pernambuco, com capacidade geral de atendimento variando entre 411 e 831 leitos. Em nosso estudo consideramos doze atributos relacionados a cada paciente (1- hospital que o paciente foi atendido; 2- idade; 3- sexo; 4- hipertensão arterial; 5- obesidade; 6- doença renal crônica; 7- doença pulmonar crônica; 8- uso de ventilação mecânica; 9- necessidade de diálise; 10- posição prona; 11- utilização de UTI; 12- custo total).

Analisamos os dados para identificar fontes de variabilidade e relações mais fortes entre os atributos e a variável custo. Ressalta-se a grande variabilidade dos custos médios por hospital, conforme mostrado na coluna 2 da Tabela 1. Também verificamos um grande impacto no custo final relacionado com o atributo 11 (utilização de UTI). Ressalta-se também a grande esparsidade do atributo idade, o que nos levou a discretizar a idade em faixas (intervalos de 20 anos), permitindo que um modelo de predição possa encontrar mais similaridades (padrões) entre os pacientes. A partir dessa caracterização, os dados foram divididos em treino e teste (50% cada), realizando a classificação com 10 repetições. A acurácia do modelo foi avaliada usando as partições de teste. Os resultados são apresentados na Tabela 1, onde cada linha referencia um determinado hospital e cada coluna é um cenário definido a seguir. O cenário 1 (coluna “Predição UTI”) corresponde a uma tarefa de predição binária para determinar se é possível prever se o paciente irá para a UTI (ou não), utilizando apenas os atributos obtidos na admissão do paciente (Atributos 2-7). Os resultados mostram a dificuldade na predição utilizando os dados de admissão hospitalar, que é um pouco melhor do que um modelo totalmente aleatório (50%).

A outra tarefa investigada é a predição de custo final (variável predita), categorizado em três faixas de forma uniforme (baixo, médio e alto). Ao realizar a categorização de custo de forma uniforme, o modelo é simplificado, pois não há desbalanceamento entre classes. Nesse cenário de predição da faixa de custo, utilizamos os atributos 2-7 como preditores e o atributo 12 como desfecho (coluna “Predição de custo”). Novamente, temos uma predição estatisticamente diferente para cada hospital. Por fim, definimos um último cenário, relacionado com a coluna “Predição custo com dados dinâmicos”, onde temos a predição da faixa de custo adicionando informações relacionados a UTI (atributos 8-10). Podemos observar que há ganho em média em todos os hospitais, alcançando melhoria de até 70% (0.34 para 0.58 no hospital 1), na acurácia média.

Tabela 1. Acurácia e intervalo de confiança com 95%. (UTI: Unidade de terapia intensiva)

Hospital	Custo Médio (R\$)	Predição UTI	Predição de custo	Predição custo com dados dinâmicos
1	6395.9	0.54 +/- 0.01	0.34 +/- 0.04	0.58 +/- 0.04
2	15767.1	0.59 +/- 0.04	0.29 +/- 0.06	0.39 +/- 0.05
3	11944.3	0.56 +/- 0.04	0.29 +/- 0.04	0.46 +/- 0.04
4	5158.6	0.58 +/- 0.03	0.40 +/- 0.04	0.41 +/- 0.06
5	13316.6	0.61 +/- 0.03	0.37 +/- 0.04	0.49 +/- 0.04

4. DESAFIOS E APRENDIZADOS

Definimos e avaliamos diferentes cenários de predição, obtendo até 61% de acurácia na tarefa de predição de UTI e até 58% de acurácia na predição de faixas de custo. Como desafios, temos a escassez de dados relativa à quantidade de pacientes e a grande variabilidade entre hospitais em relação aos custos.

5. REFERÊNCIAS

1. Cardoso, Ricardo B., et al. "Comparison of COVID-19 hospitalization costs across care pathways: a patient-level time-driven activity-based costing analysis in a Brazilian hospital." BMC H. S. Res. 23.1 (2023).



Proposição de uma ontologia de domínio para pacientes oncopediátricos

M. C. Pereira^{1,2,3}, M. Sinigaglia^{1,3}, S. C. Cazella²

¹ Instituto do Câncer Infantil, Porto Alegre, RS

² Programa de Pós-graduação em Ciências da Saúde, Universidade Federal de Ciências da Saúde de Porto Alegre, Porto Alegre, RS

³ Instituto Nacional de Ciência e Tecnologia em Biologia do Câncer Infantil e Oncologia Pediátrica - INCT BioOncoPed, Porto Alegre, RS

marianacp@ufcspa.edu.br, msinigaglia@ici.org, silvioc@ufcspa.edu.br

Resumo. O incremento na geração de dados criou a necessidade das instituições de saúde desenvolverem bases de dados consistentes e padronizadas para o acompanhamento dos pacientes e gestão eficiente do conhecimento. Todavia, a dispersão da informação ainda é um desafio a ser enfrentado e as ontologias têm sido utilizadas para sanar esta problemática. Este resumo apresenta uma pesquisa em andamento que visa à proposição de uma ontologia de domínio para a visualização sistêmica do paciente oncopediátrico. O conhecimento será adquirido por meio da triagem de fichas clínicas, protocolos clínicos e diretrizes terapêuticas, com a participação de profissionais da área. A ontologia será estruturada via ferramenta Protégé e será utilizada a Metodologia 101, podendo-se adicionar outros métodos de modelagem ontológica. Espera-se propor uma ontologia de domínio que represente o conhecimento relacionado aos pacientes oncopediátricos, facilitando o entendimento de como as ontologias de domínio podem contribuir na representação e manipulação dos dados em saúde neste contexto.

Palavras-chave: Interoperabilidade de informação em saúde; Oncologia pediátrica; Ontologia.

Abstract. The increase in data generation created the need for health institutions to develop consistent and standardized databases for the monitoring of patients and for the efficient management of knowledge. However, the dispersion of information is still a challenge to be faced and ontologies have been used to remedy this problem. This paper presents a research in progress that aims to propose a domain ontology for the systemic view of oncopediatric patients. The knowledge will be acquired through the screening of clinical records, clinical protocols and therapeutic guidelines, with the participation of professionals in the area. The ontology will be structured using the Protégé tool and Methodology 101 will be used, with the possibility of adding other ontological modeling methods. It is expected to propose a domain ontology that represents the knowledge related to oncopediatric patients, facilitating the understanding of how domain ontologies can contribute to the representation and manipulation of health data in this context.

Keywords: Health Information Interoperability, Ontology, Pediatric Oncology.

Nome do projeto: Proposta de ontologia de domínio de pacientes oncopediátricos visando à aplicação em sistemas inteligentes

1. CONTEXTO E MOTIVAÇÃO

Os tumores pediátricos representam 3% dos tumores malignos mundiais e são a primeira causa de morte por doença em crianças e adolescentes no Brasil (1). Tal realidade é desoladora, visto que 80% dos pacientes acometidos pela doença poderiam ser curados se diagnosticados precocemente e tratados em centros de referência (1). Além disso, os serviços de saúde enfrentam os desafios decorrentes da dispersão de registros clínicos, que impactam na eficiência dos profissionais, no acompanhamento clínico e na padronização e integração das informações (2). Assim, a proposição de uma ontologia de domínio se torna crucial para a gestão de dados e do conhecimento, principalmente pelo seu impacto na qualidade dos registros em saúde e na capacidade da tomada de decisão baseada em evidências para a prestação da assistência em saúde.

1.1 Objetivo

Propor uma ontologia de domínio de pacientes oncopediátricos atendidos no Instituto do Câncer Infantil visando à aplicação em Sistemas Inteligentes.



2. ATIVIDADES PRINCIPAIS

As atividades principais realizadas para execução deste projeto são as seguintes: 1) Mapear e pré-processar variáveis e suas possibilidades de entrada em fichas clínicas de um estudo epidemiológico, bem como desenvolver um dicionário de dados para garantir a padronização das informações com base em parâmetros internacionais; 2) Descrever o fluxo, objetivos e hipóteses de um banco de dados para estudos de acompanhamento de pacientes oncopediátricos e especificar as relações terminológicas e conceituais entre os domínios oncopediátricos; 3) Buscar por ontologias para reuso e consultar especialistas e fontes consolidadas no domínio da oncopediatria.

3. DESENVOLVIMENTO DO TRABALHO

Para o mapeamento e pré-processamento de variáveis e suas possibilidades de entrada, serão utilizadas as fichas clínicas do estudo “Acompanhar para transformar: um olhar integrado para o câncer infantojuvenil a longo prazo, considerando aspectos clínicos, psicológicos e sociais”, CAAE 52044221.8.1001.5327, disponibilizadas pelo Instituto do Câncer Infantil. As fichas não terão quaisquer dados de pacientes e, em razão de não envolver seres humanos, o presente projeto não necessita de apreciação em Comitê de Ética Pesquisa. O projeto foi registrado e aprovado pela Comissão de Pesquisa da Universidade Federal de Ciências da Saúde de Porto Alegre sob número 201/2022. A ontologia será construída por meio da ferramenta Protégé e modelada via Metodologia 101 (3), entretanto, poderá ser necessário combinar outros métodos para estruturar a ontologia conforme sua complexidade adquirida.

4. DESAFIOS E APRENDIZADOS

Os desafios referem-se à identificação dos métodos para a modelagem da ontologia focada na área da saúde e a seleção do método mais alinhado com a finalidade desta ontologia. Destaca-se a dificuldade de identificar ontologias disponíveis que se alinhem às mudanças de protocolos clínicos e diretrizes terapêuticas. A título de exemplo, através de pesquisas científicas, o conhecimento relacionado aos tumores é aprimorado e, portanto, um determinado tumor poderia ter uma classificação histológica e, anos depois, esta classificação é modificada, o que exige a atualização do conhecimento. Caso esse conhecimento esteja representando em ontologias, faz-se necessário atualizá-las constantemente. Os aprendizados se relacionam à evolução do olhar às variáveis inseridas na ontologia, sobretudo para não limitar a sua abrangência e reuso. Ademais, a caracterização abrangente dos pacientes oncopediátricos é fundamental para as intervenções clínicas, proposição de políticas de saúde e padronização da informação para fins de gestão e ensino em saúde, sendo estes os principais aprendizados decorrentes do estudo.

5. AGRADECIMENTOS

A equipe do projeto agradece à Fundação de Amparo à Pesquisa do Rio Grande do Sul pelo apoio financeiro através do Projeto FAPERGS 22/2551-0000390-7 (RITE CIARS).

6. REFERÊNCIAS

1. INCA. Instituto Nacional de Câncer. Câncer infantojuvenil [Internet]. (2023). Disponível em: <https://www.inca.gov.br/tipos-de-cancer/cancer-infantojuvenil>. Acessado em 06/08/2023.
2. Teixeira Livia M. D, Almeida, Maurício B. “Aspectos ontológicos e epistemológicos em terminologias clínicas: em busca de interoperabilidade semântica no ambiente clínico”. Encontros Bibli [Internet]. 24(55):1-21 (2019). Disponível em: <https://periodicos.ufsc.br/index.php/eb/article/view/1518-2924.2019.e57996>. Acessado em 03/01/2023.
3. Noy, Natalya F., McGuinness Deborah L. “Ontology Development 101: A Guide to Creating Your First Ontology”. [Internet]. Disponível em: https://protege.stanford.edu/publications/ontology_development/ontology101.pdf. Acessado em 15/01/2023.



Segmentação semântica de ECG por meio de técnicas de self-learning

Matheus F. A. A. Oliveira (aluno)¹, Derick M. Oliveira (aluno)¹, Wagner Meira Jr. (orientador)¹

¹Universidade Federal de Minas Gerais, Departamento de Ciência da Computação

akira.matheus@dcc.ufmg.br, derimath@dcc.ufmg.br, meira@dcc.ufmg.br

Resumo. Aplicações de aprendizado de máquina para medicina são muito importantes, em especial para a cardiologia. A segmentação de um ECG é uma tarefa importante, mas até então não solucionada. Nesse trabalho propomos uma abordagem que consegue utilizar informação de ECGs não rotulados para a segmentação. Como resultado parcial, criamos um modelo de self-learning que é capaz de reconstruir o sinal com precisão.

Abstract. Machine learning applications for medicine are very important, especially for cardiology. The segmentation of an ECG is an important, however, unresolved task. In this work, we propose an approach that manages to use information from unlabeled ECGs for segmentation. As a partial result, we have created a self-learning model capable of accurately reconstructing the signal.

Palavras-chave: Segmentação semântica; Aprendizado de Máquina; Cardiologia.

1. CONTEXTO E MOTIVAÇÃO

Diversos trabalhos propondo modelos de ML para auxiliar a prática médica foram propostos nos últimos anos com técnicas que melhoram o desempenho de trabalhos rotineiros com extrema precisão. Tarefas que até então necessitavam de especialistas com alto grau de conhecimento na área.

A cardiologia é uma área de grande importância para aplicação de modelos de machine learning, dado ao grande volume de dado e por se tratar de uma vital área médica para salvar vidas. É estimado que 30% da população global e do Brasil morrem por doenças cardiovasculares [1]. Uma tarefa de grande importância para os cardiologistas é a segmentação do exame de Eletrocardiograma (ECG) nos segmentos usados na prática médica, que até então são fornecidos por algoritmos proprietários fechados com baixa precisão, na população brasileira que é muito diversa.

A segmentação do ECG é uma tarefa custosa, então, apesar do grande volume de ECG feitos diariamente, poucos exames são segmentados para aplicação de técnicas de machine learning recentes em tarefas supervisionadas.

1.1 Objetivo

Esse trabalho tem como objetivo gerar um modelo de aprendizado semi-supervisionado, aprendendo a representação simples de um ECG por meio de uma tarefa não supervisionada e posteriormente adaptada por finetuning com a tarefa de segmentação semântica.

Utilizamos uma técnica de self-learning conhecida como Variational AutoEncoder (VAE) [2] para extrair a representação latente do ECGs e pretendemos utilizar uma rede convolucional, estado-da-arte em classificação para ECGs [3], na tarefa de segmentação semântica para identificar as ondas.

Como resultados parciais, conseguimos extrair as variáveis latentes do ECG com o MSE de 0,0004, sendo um valor bem baixo para a tarefa. O próximo passo é adaptação para a tarefa de segmentação semântica.

2. ATIVIDADES PRINCIPAIS

As atividades desenvolvidas, foram:

- Pre-processamento sobre base de dados dos ECGs para todos serem normalizados
- Implementação de uma rede neural VAE para self-learning

Os próximos passos



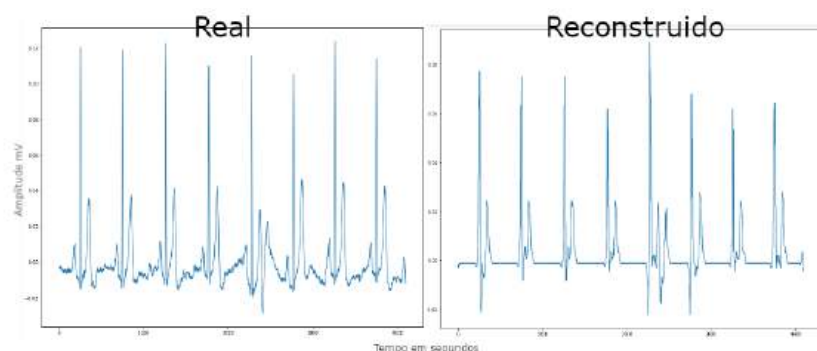
- Implementação de rede convolucional para segmentação semântica

3. DESENVOLVIMENTO DO TRABALHO

Para o desenvolvimento desse trabalho foram coletados 2.356.873 exames eletro-cardiológicos de 1.621.538 pacientes de 811 municípios do estado de Minas Gerais, coletados entre 2010 e 2016. Os dados foram obtidos em parceria com a Rede de Telessaúde de Minas Gerais, coordenada pelo Centro de Telessaúde do Hospital das Clínicas da Universidade Federal de Minas Gerais (CTHC-UFMG). O ECG obtido são sinais temporais de 10 segundos com 12 derivações.

O modelo de self-learning aplicado foi um VAE, adaptado para dados 1D que tem duas regiões de aprendizado (1) Encoder, cujo objetivo é gerar o espaço latente utilizando como entrada o sinal do ECG; (2) Decoder, com a entrada o espaço latente gerado pelo encoder e como saída uma reconstrução do sinal original. A tarefa de reconstrução é minimizar o sinal reconstruído com o sinal de entrada com uma função de perda erro médio quadrático (MSE).

Como resultados temos um modelo com o MSE de 0,0004, e um exemplo de reconstrução na Figura ao lado como resultado qualitativo. Como podemos ver, a reconstrução é bem próxima do sinal original a menos de



ruído e pequenas deflexões no sinal. A segmentação do ECG ainda é uma tarefa em aberto dado a dificuldade de gerar labels para essa tarefa. Novas técnicas em aprendizado self-supervised podem melhorar a segmentação desses diagnósticos refletindo em diagnósticos mais acurados o que pode melhorar a qualidade de vida dos pacientes.

4. DESAFIOS E APRENDIZADOS

Os principais desafios para esse projeto são:

- Interdisciplinaridade com área da medicina cardiovascular;
- Adaptação de técnicas para séries temporais;
- Implementação de modelos estado-da-arte em segmentação semântica e self-learning;
- Otimização de parâmetros para adequar as técnicas usadas.

5. REFERÊNCIAS

1. Oliveira GM, Brant LC, Polanczyk CA, Biolo A, Nascimento BR, Malta DC, Souza MD, Soares GP, Xavier Junior GF, Machline-Carrion MJ, Bittencourt MS. Estatística Cardiovascular–Brasil 2020. Arquivos brasileiros de Cardiologia. 2020 Sep 28;115:308-439.
2. Doersch C. Tutorial on variational autoencoders. arXiv preprint arXiv:1606.05908. 2016 Jun 19.
3. Ribeiro AH, Ribeiro MH, Paixão GM, Oliveira DM, Gomes PR, Canazart JA, Ferreira MP, Andersson CR, Macfarlane PW, Meira Jr W, Schön TB. Automatic diagnosis of the 12-lead ECG using a deep neural network. Nature communications. 2020 Apr 9;11(1):1760.



Sign Language Recognition System for Deaf Patients

Lucca Fagundes Ramos de Oliveira (aluno)¹, Lucas Rocha Valle (aluno)², Luiza Marinho Motta Santa Rosa (aluno)³, Raquel Prates (orientador)⁴, Zilma Reis (orientador)², Mário Fernando Montenegro Campos (orientador)⁴, Milena Soriano Marcolino (orientador)²

¹Faculdade de Medicina, Universidade Federal de Ouro Preto, Ouro Preto, MG

²Faculdade de Medicina, Universidade Federal de Minas Gerais (UFMG), Belo Horizonte, MG

³Faculdade Ciências Médicas de Minas Gerais, Belo Horizonte, MG

⁴Departamento de Ciências da Computação, UFMG, Belo Horizonte, MG

fagundes.lucca8@gmail.com, lucasrochavalle@gmail.com, luiza.motta26@gmail.com,
rprates@dcc.ufmg.br, zilma@ufmg.br, mario@dcc.ufmg.br, milenamarc@gmail.com

Abstract. *Hearing-impaired individuals are more likely to suffer from preventable health issues, including iatrogenic events, adverse problems in acute care and chronic diseases and, for this reason, it becomes evident that there are important barriers related to communication between healthcare workers and the deaf individuals. Considering these facts, the elaboration of a systematic review was initiated with the aim to gather and evaluate the applicability and usability of the different communication systems developed for deaf people. A literature search was conducted in order to select articles that described the creation of communication systems for hearing-impaired people, applied in the health context, which were tested somehow on human users. After the literature search, 10.890 articles were obtained, of which 1.157 duplicate articles were removed through Rayyan, leaving 9.733 articles. In the phase of screening by title and abstract, 9.632 articles were excluded, leaving 101 articles for full text reading. Currently, the selection of these articles is still in progress.*

Keywords: Artificial Intelligence; Sign Language; Hearing Loss.

Project Name: Captar-Libras:Sistema de Comunicação por vídeos para surdos aplicado ao pré-atendimento médico (MFMC)

1. CONTEXT AND MOTIVATION

According to World Health Organisation (WHO), approximately 430 million people, which accounts for over 5% of the world's population, have some disabling hearing loss. It is estimated that by 2050 over 700 million people will have some hearing impairment.¹ Communication issues between healthcare workers and the deaf community have a huge impact on these individuals' health, making them more susceptible to suffer from preventable health issues, including adverse problems in acute care and chronic diseases². Within this context, several technologies³, using artificial intelligence, have been developed, all around the world, in order to facilitate and promote better interaction with the deaf community.

1.1 Objective

This study aims to systematically review the evidence on human-computer approaches for establishing or assisting deaf patients during healthcare encounters, which were tested in human users.

2. KEY ACTIVITIES

Initially, the undergraduate students, together with the supervisors, worked in the systematic review protocol⁴ and established the search strategy. Subsequently, the undergraduate students conducted the search and, using the Rayyan application, removed the duplicated articles and selected the studies of interest, considering the inclusion criteria previously established.



3. WORK DEVELOPMENT

A literature search, with no language restriction, was performed using Web of Science, Medline (assessed through Pubmed), ACM and Instituto de Engenheiros Eletricistas e Eletrônicos (IEEE) Xplore, without using any chronological filters, to identify relevant studies. To support the research in Pubmed, combinations of the MeSH terms "Artificial Intelligence", "Communication Aids for Disabled", "Sign Language" and "Hearing Loss" were used in the advanced search bars. After the literature search, 10.890 articles were obtained, of which 1.157 duplicate articles were removed through Rayyan, leaving 9.733 articles. Studies that presented the description of a sign language recognition system tested on human users or their videos, which could be applied in a health context, were eligible. The exclusion criteria were basically the development of the sign language communication system for other contexts than health care and systems without testing in human users. Titles and abstracts were independently screened by two investigators and any disagreement was resolved by a third reviewer. The articles considered potentially relevant were posteriorly selected for a full-text reading. In the phase of screening by title and abstract, 9.632 articles were excluded, leaving 101 articles for full text reading. This study is still in progress. The investigators are currently reading the complete articles to check for eligibility.

4. CHALLENGES AND LEARNINGS

Despite the great number of articles found in literature search, few studies described sign language recognition systems tested in human users and even fewer used in healthcare contexts. This finding shows that sign language recognition systems are not being widely explored by healthcare professionals and, despite technological advances, artificial intelligence's practical application in this context is still limited. The potential of such systems to improve the inclusion of people with disabilities and their social functioning is undeniable, however there is a long path until these benefits are achieved.

This study emphasizes the importance of scientific research investments, as it significantly accelerates project development and increases the resources in the field of human-machine interaction.

5. REFERENCES

1. World Health Organization. Deafness and Hearing Loss [Internet]. Available at: <https://www.who.int/news-room/fact-sheets/detail/deafness-and-hearing-loss> (31 July 2023, date last accessed).
2. Bartlett G, Blais R, Tamblyn R, Clermont RJ, MacGibbon B. Impact of patient communication problems on the risk of preventable adverse events in acute care settings. *CMAJ*. 2008 Jun 3;178(12):1555-62. doi: 10.1503/cmaj.070690. PMID: 18519903; PMCID: PMC2396356.
3. Papastratis I, Chatzikonstantinou C, Konstantinidis D, Dimitropoulos K, Daras P. Artificial Intelligence Technologies for Sign Language. *Sensors (Basel)*. 2021 Aug 30;21(17):5843. doi: 10.3390/s21175843. PMID: 34502733; PMCID: PMC8434597.
4. <https://doi.org/10.17605/OSF.IO/FPEMR>



small RNA Metavir: uma ferramenta baseada em aprendizado de máquina para identificação de vírus em amostras complexas independente de análises de similaridade de sequência

Hebert N. A. Costa¹, João Paulo P. de Almeida¹, Bruno M. Silva², Paulo R. F. Júnior³, Ulisses B. Corrêa³, Mauro M. Teixeira¹, Eric R. G. R. Aguiar⁴, João T. Marques (orientador)¹

¹*Instituto de Ciências Biológicas da UFMG, Belo Horizonte, MG*

²*Instituto Tecnológico Vale - Belém, PA*

³*Hub de Inovação em Inteligência Artificial CDTEC UFPEL, Pelotas, RS*

⁴*Centro de Biotecnologia e Genética da UESC, Ilhéus, BA*

hebert.costa.jhale@gmail.com, joaopaulobio@ufmg.br, bruno.marques@pq.itv.org
paulo.ferreira@inf.ufpel.edu.br, ulisses@inf.ufpel.edu.br, mmtex.ufmg@gmail.com, ergraguiar@uesc.br,
jtm@icb.ufmg.br

Resumo. A descoberta de novos vírus e o monitoramento eficiente da sua circulação tiveram um grande impacto com o uso de abordagens metagenômicas a partir de dados de sequenciamento de nova geração. Entretanto, abordagens de metagenômica para a identificação de sequências virais em animais possuem limitações por depender de similaridade com referências já conhecidas e de não conseguir distinguir elementos virais endógenos, chamados de EVEs (falsos positivos de vírus circulantes). Apresentamos uma ferramenta computacional para detecção e classificação de sequências virais utilizando dados de sequenciamento de pequenos RNAs produzidos pela resposta antiviral do hospedeiro. A ferramenta consiste num pipeline automatizado para montar e produzir representações estruturadas para sequências com base em atributos numéricos de perfis de tamanho de pequenos RNAs. A ferramenta possui um classificador ensemble para distinção entre sequências virais exógenas e EVEs (falsos positivos), treinado com uma base de sequências virais curadas geradas a partir do processamento de 122 bibliotecas de pequenos RNAs de mosquitos coletados ao redor do planeta. A aplicação de modelos de aprendizado não-supervisionado atesta a elevada capacidade distintiva e coerência espacial das representações geradas pela ferramenta. Testes do modelo ensemble em bibliotecas públicas de pequenos RNAs de mosquitos evidenciam alto desempenho do classificador na distinção de sequências virais de EVEs. Nossa ferramenta computacional desenvolvida oferece um recurso valioso para aprimorar futuros estudos metagenômicos de vetores de doenças de grande relevância para saúde pública.

Palavras-chave: Pequenos RNAs; Engenharia de Dados; Mosquito, Metagenômica.

Nome do projeto: Desenvolvimento de uma ferramenta baseada em aprendizado de máquina independente de similaridade de sequências para identificação de vírus a partir de dados de metagenômica

1. CONTEXTO E MOTIVAÇÃO

Vírus são parasitas intracelulares que, devido a suas altas taxas de mutação e replicação, podem rapidamente se adaptar, infectar novos hospedeiros e adquirir patogenicidade. Como exemplo, os arbovírus - vírus transmitidos de artrópodes (como mosquitos) para vertebrados - representam uma grande ameaça à saúde humana devido aos grandes impactos socioeconômicos causados por arboviroses como a dengue. Na análise do viroma de mosquitos e outros organismos podem ser encontrados, além de uma grande variedade de vírus, elementos endógenos virais (EVEs). EVEs são resquícios da integração viral ao genoma do hospedeiro (i.e. não são mais vírus) que podem distorcer estudos de viroma caso não sejam corretamente distinguidas de sequências virais reais (exógenas). Como ainda não há vacinas ou antivirais eficientes contra a maioria dos vírus, o monitoramento epidemiológico (que requer a identificação de vírus em meio a amostras de tecidos



complexos), é fundamental para controlar e combater novas epidemias virais. Métodos computacionais de identificação de vírus a partir de dados de sequenciamento de RNA geralmente se baseiam em similaridade de sequências com referências previamente conhecidas. Por essa razão estudos frequentemente resultam em grandes quantidades de sequências não caracterizadas que podem conter novas sequências virais. Essa limitação demanda alternativas de métodos que sejam independentes de similaridade de sequência, a exemplo dos que exploram perfis de tamanhos de pequenos RNAs. A utilização desses atributos possibilita a representação computacional de sequências de pequenos RNAs com alta coerência e que pode ser utilizada em aplicações de IA para descoberta de novos vírus e vigilância de patógenos virais.

1.1 Objetivo

Oferecer à comunidade científica uma ferramenta computacional que explora perfis de tamanho de pequenos RNAs para representação estruturada e classificação automática de sequências virais com uso de aprendizado de máquina. A ferramenta inclui um pipeline para engenharia de dados brutos de sequências metagenômicas e um classificador ensemble treinado para distinção automatizada de sequências metagenômicas entre virais exógenas ou EVEs. Todas as funcionalidades são encapsuladas numa imagem linux que permite utilização simplificada via interface de linha de comando e que poderá ser integrada a sistemas de monitoramento de vírus, prevenção de epidemias e em aplicações de IA para a área de saúde.

2. ATIVIDADES PRINCIPAIS

Treinamento e seleção de modelos classificadores utilizados para diferenciação de sequências exógenas e EVEs; encapsulamento da aplicação junto a suas dependências; realização de testes; consolidação e análise de dados de resultados dos testes; versionamento e padronização de releases do repositório público; planejamento de melhorias e atualizações do sistema;

3. DESENVOLVIMENTO DO TRABALHO

O desenvolvimento e testes do pipeline de engenharia de dados, bem como do encapsulamento da aplicação numa imagem linux, foi realizado em servidores linux acessados via protocolo SSH. Os testes de execução de contêineres foram feitos com uso das runtimes Docker e Podman, havendo sucesso com ambas. Quatro máquinas, com diferentes configurações de hardware, foram utilizadas para testes onde se verificou variações de tempo de execução. A eficiência na geração de perfis de tamanho de pequenos RNAs foi atestada através de inspeção visual de gráficos de perfil de tamanhos gerados automaticamente pelo pipeline. A qualidade da representação estruturada de sequências metagenômicas de pequenos RNAs gerada foi validada com uso de algoritmos de aprendizado não supervisionado onde se verificou coerência espacial e agrupamento hierárquico adequados. A eficiência do classificador na diferenciação de sequências virais endógenas e EVEs foi atestada via processo de validação cruzada de diferentes modelos ensemble e avaliação de curvas ROC com dataset proveniente de curadoria manual. Os melhores modelos com os algoritmos catboost e Random Forest apresentaram valores AUC maiores que 0.9. A primeira versão da ferramenta já se encontra disponível para download e uso em sistemas Unix: <https://github.com/rnai-bioinfo/small-rna-metavir>.

4. DESAFIOS E APRENDIZADOS

A natureza multidisciplinar da equipe trouxe o desafio da integração entre domínios das ciências biológicas, imunologia, virologia, bioinformática, aprendizado de máquina e engenharia de sistemas. Entre os maiores aprendizados destacam-se, além dos conhecimentos em biologia, o aumento de habilidades com treinamento e seleção de modelos de aprendizado de máquina, terminal linux, ferramentas de containerização.

5. REFERÊNCIAS

1. Aguiar ER, Olmo RP, Paro S, Ferreira FV, de Faria IJ, Todjro YM, Lobo FP, Kroon EG, Meignin C, Gatherer D, Imler JL. Sequence-independent characterization of viruses based on the pattern of viral small RNAs produced by the host. Nucleic acids research. 2015 Jul 27;43(13):6191-206.



Somatório de curvas de Kaplan-Meier: um equivalente meta-analítico para sobrevida global que pode se beneficiar significativamente do uso de inteligência artificial.

Matheus Henrique Leite e Silva ¹, Daniel Simões Reis dos Santos ¹, Davi Bastos Wolff ¹

¹Faculdade de Medicina - Universidade Federal de Minas Gerais - UFMG, Belo Horizonte, MG

m4henqk@ufmg.br, danielsrs@ufmg.br, dwolff@ufmg.br

Resumo. Medicina baseada em evidências é o atual padrão-ouro para tomada de decisões em saúde, sendo regida em grande parte por revisões sistemáticas, acompanhadas ou não por meta-análises, que transformam o corpo de pesquisa original em resultados sintéticos e orientados à prática clínica. Muito esforço foi empenhado para otimização dos formatos empregados na publicação de dados primários e, da mesma forma, mecanismos para relatar achados revisionais também são alvo de intensa investigação. As curvas de Kaplan-Meier são a representação gráfica de escolha para comparação de probabilidade de sobrevida global entre grupos em ensaios clínicos randomizados, contudo, seu equivalente para revisões sistemáticas ainda está no início de seu desenvolvimento e popularização: curvas de Kaplan-Meier multi-ensaio. O somatório de dados de sobrevida de múltiplos ensaios clínicos confere maior acurácia aos resultados, de forma similar às meta-análises, e permite potencial aumento do nível de evidência para esse parâmetro clínico. Interpõem-se, contudo, questões técnicas como escolha de softwares e variabilidade interobservador. O presente resumo busca relatar uma das possibilidades de execução desse tipo de análise estatística e aplicá-la em um cenário de vida-real.

Palavras-chave: Estimativa de Kaplan-Meier; Lutécio; Inteligência Artificial.

1. CONTEXTO E MOTIVAÇÃO

Curvas de Kaplan-Meier são utilizadas desde a década de 1950 (1), sendo especialmente úteis em contextos de incompletude de dados, para reportar a probabilidade de um sujeito sobreviver até certo ponto no tempo. Essa representação gráfica é uma estimativa não-paramétrica, sem pressuposições sobre a distribuição das variáveis, podendo ser livremente incluída de acordo com o CONSORT Statement para ensaios clínicos randomizados (2). A Sobrevida Global e seus equivalentes válidos são desfechos primários amplamente utilizados para comparação de tratamentos e outros fatores prognósticos, com implicações diretas para decisões de pacientes e de médicos assistentes. Evidentemente, tanto o rigor metodológico quanto a forma de apresentação escolhida são vitais para a garantia da confiabilidade desse desfecho clínico, o que justifica a busca por padronização com objetivo de aumentar a qualidade das recomendações nele fundamentadas.

1.1 Objetivo

Divulgar o método IPDfromKM, em Shiny (3), que são pacotes das linguagens de programação web app e R, respectivamente, para agregação de dados e reconstrução de curvas de Kaplan-Meier, ainda não efetivamente difundido no cenário brasileiro, mas que pode se tornar uma ferramenta de auxílio na produção de dados revisionais comparável às meta-análises; relatar as experiências e conhecimentos adquiridos pelos autores.

2. ATIVIDADES PRINCIPAIS

Foi explorada uma metodologia de extração, tratamento e reconstrução de dados com finalidade de controle de qualidade da técnica e familiarização através da comparação de resultados independentes, mas padronizados, de dois dos autores, MS e DS. As curvas de Kaplan-Meier foram obtidas, em alta resolução, diretamente do artigo (4). O software ScanIt (5) foi empregado para determinação das coordenadas da curva, através do modo *Trace curve* com *window size* de 8, seguida por correção manual. Subsequentemente, as curvas de Kaplan-Meier foram reconstruídas com uso da aplicação IPDfromKM.



3. DESENVOLVIMENTO DO TRABALHO

Na investigação da forma ideal de relatar desfechos estudados na literatura primária, foi aplicado o método IPDfromKM (3) para reconstrução de dados de pacientes individuais expressos em curvas de Kaplan-Meier. Dessa forma, seria gerada maior familiaridade com a metodologia e validação da técnica para uso em um contexto regional. Foi escolhido o ensaio clínico randomizado TheraP (4), estudo seminal para o início do processo de aprovação do radionuclídeo ^{177}Lu -PSMA-617 pela FDA, que envolveu também outros ensaios clínicos e ocorreu em março de 2022. Os resultados preliminares dessa abordagem foram a elaboração de um fluxograma com as etapas envolvidas na reconstrução de uma curva de Kaplan-Meier, validada pela comparação à curva original, e que deve ser subsequentemente somada a outras curvas de estudos semelhantes, de forma meta-analítica, com importantes aplicações na avaliação de tecnologias em saúde.

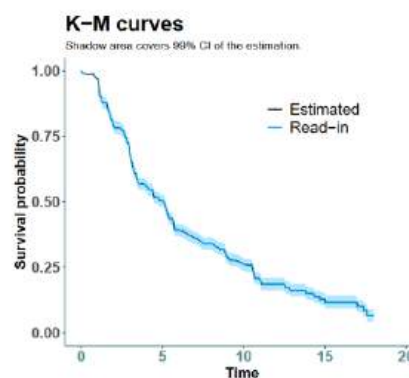


Figura1: Fluxograma do processo de reconstrução dos dados individuais e curva resultante

4. DESAFIOS E APRENDIZADOS

No processo de aplicação do IPDfromKM, os desafios iniciais foram a compreensão do funcionamento de sua IA, o maior domínio com a linguagem R e o entendimento do uso da sobrevida global como desfecho primário. Havia, ainda, a necessidade de escolha de um software para conversão cartesiana, o que levou à procura e identificação do ScanIt, dada a possibilidade de automatização, disponibilidade de ferramentas de correção manual e licença gratuita. Durante seu uso, um dos maiores aprendizados foi a exploração de suas ferramentas na busca por um processo padronizado, visando a uma conversão reprodutível.

5. REFERÊNCIAS

1. Rich JT, Neely JG, Paniello RC, Voelker CCJ, Nussenbaum B, Wang EW. A practical guide to understanding Kaplan-Meier curves. *Otolaryngology–Head and Neck Surgery*. 2010 Sep;143(3):331–6.
2. Butcher NJ, Monsour A, Mew EJ, Chan AW, Moher D, Mayo-Wilson E, et al. Guidelines for Reporting Outcomes in Trial Reports. *JAMA*. 2022 Dec 13;328(22):2252.
3. Liu N, Zhou Y, Lee JJ. IPDfromKM: reconstruct individual patient data from published Kaplan-Meier survival curves. *BMC Medical Research Methodology*. 2021 Jun 1;21(1)
4. Hofman MS, Emmett L, Sandhu S, Iravani A, Joshua AM, Goh JC, et al. [^{177}Lu]Lu-PSMA-617 versus cabazitaxel in patients with metastatic castration-resistant prostate cancer (TheraP): a randomised, open-label, phase 2 trial. *The Lancet* [Internet]. 2021 Feb;397(10276):797–804. Available from: https://www.petermac.org/sites/default/files/media-uploads/2020%20Lancet%20TheraP%20-%20Hofman_0.pdf
5. AmsterdamCHEM. ScanIt [Computer software]. Version 2.0.8.0. AmsterCHEM; 2021



Transformer generativo auto-regressivo para detecção de anomalias em imagens de ressonância magnética de cérebros

Pedro Dutenhofner¹, Turi Rezende¹, Diogo Chaves¹, Derick Oliveira¹, Gisele Pappa¹, Wagner Meira Jr.¹

¹Universidade Federal de Minas Gerais, Belo Horizonte, MG

pedroroblesduten@ufmg.br, turi@ufmg.dcc.br, diogo.tuler@gmail.com, derickmath@dcc.ufmg.br,
glpappa@dcc.ufmg.br, meira@ufmg.dcc.br

Resumo. A detecção automática de doenças em imagens médicas é algo de grande valor para a sociedade, entretanto, devido a escassez de dados disponíveis com rótulos, este processo se torna bastante limitado. Neste sentido, este artigo apresenta um método generativo “self-supervised” para detecção de anomalias em imagens MR de cérebros baseado em um transformer generativo auto-regressivo. Além de explicações teóricas e detalhes dos modelos, apresentamos resultados quantitativos e qualitativos.

Abstract. The automatic detection of diseases in medical images is something of great value to society, however, due to the scarcity of data available with label, this process becomes quite limited. In this sense, this article presents a self-supervised generative method for detecting anomalies in brain MR images based on an auto-regressive generative transformer. In addition to theoretical explanations and model details, we present quantitative and qualitative results.

Palavras-chave: Transformer; Anomalia; Generativo; Cérebros.

Nome do projeto: IA em Saúde

1. CONTEXTO E MOTIVAÇÃO

A detecção de doenças em imagens de Ressonância Magnética (MRI) é uma ferramenta vital no diagnóstico precoce e na gestão de condições médicas. Contudo, dados médicos, ainda que abundantes em certas situações, possuem um baixo percentual de rotulação, fato que dificulta o uso de modelos supervisionados. Neste trabalho, apresentamos um modelo generativo “self-supervised” para detecção de anomalias em imagens de MRI, que se beneficia da natureza incerta de dados médicos por não precisar de rótulos.

1.1 Objetivo

Implementamos um modelo generativo baseado em um autoencoder quantizado (VQ-VAE) [1] e um transformer auto-regressivo (decoder-only) [3], associando *local-relation* do “inductive bias” de convoluções e *global-relation* pelo attention de transformer [3]. Seguimos uma abordagem apresentada em [2]. Nos baseamos na ideia de que o modelo generativo aprenda a distribuição de probabilidade de cérebros normais, a fim de detectar anomalias na estrutura cerebral, como pequenas lesões e tumores.

2. ATIVIDADES PRINCIPAIS

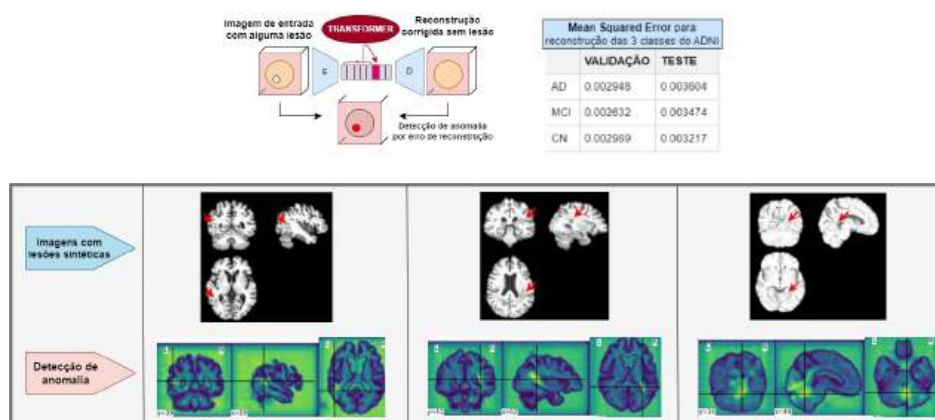
As atividades desenvolvidas envolveram processos de: pré-processamento das imagens de cérebros, implementação e treinamento de modelo de autoencoder quantizado (VQ-VAE), implementação e treinamento de modelo de transformer generativo (decoder-only) e avaliação de resultados;

3. DESENVOLVIMENTO DO TRABALHO (METODOLOGIA E RESULTADOS)

Para o desenvolvimento deste trabalho, foi utilizado o dataset público Alzheimer’s Disease Neuroimaging Initiative (ADNI) database (adni.loni.usc.edu) (5). Primeiramente foi realizado um pré-processamento envolvendo alinhamento e brain-extraction. O processo de detecção de anomalia se baseia em uma tarefa de



predição de próximo token no espaço latente de um VQ-VAE. Nesse sentido, utilizamos um autoencoder de camadas residuais convolucionais para projetar o dado de entrada para um espaço comprimido e discretizamos esse espaço latente em tokens, como proposto por [1]. A partir dos tokens obtidos, podemos treinar um transformer autoregressivo apenas com samples sem diagnóstico associado, a fim de permitir que tal modelo aprenda a distribuição de cérebros normais. No momento de inference, ao passarmos um cérebro não-normal (com alguma lesão ou tumor), o transformer generativo erraria na previsão de tokens anômalos, por estarem fora da distribuição esperada. Desse modo, caso a previsão do próximo token esteja abaixo de um determinado threshold arbitrário, substituímos o token anômalo pelo token de maior probabilidade a partir do output do transformer. Enviamos para o decoder um “espaço latente corrigido”, que será utilizado para reconstrução. Comparando o dado original com a reconstrução corrigida, podemos saber onde o transformer detectou tokens anômalos e onde estão as lesões deste determinado cérebro. Para avaliar os resultados de forma controlada, inserimos pequenos quadrados de vários tamanhos e intensidades nos dados de treino, simulando a detecção de lesões. Exemplos qualitativos podem ser observados abaixo.



4. DESAFIOS E APRENDIZADOS

Apesar de não necessitarmos explicitamente de dados rotulados, o maior desafio enfrentado durante o processo de pesquisa ocorreu em decorrência da quantidade insuficiente de dados para treinamento do transformer generativo descrito nas seções anteriores. A baixa quantidade de imagens de MRI disponíveis no dataset público ADNI causou sobreajuste no aprendizado e dificultou a generalização do modelo. Em busca de mitigar o efeito descrito, foram aplicadas nos dados técnicas de *data augmentation*.

5. REFERÊNCIAS

1. Van Den Oord, Aaron, and Oriol Vinyals. "Neural discrete representation learning." *Advances in neural information processing systems* 30 (2017).
2. Pinaya, Walter HL, et al. "Unsupervised brain imaging 3D anomaly detection and segmentation with transformers." *Medical Image Analysis* 79 (2022): 102475.
3. Esser, Patrick, Robin Rombach, and Bjorn Ommer. "Taming transformers for high-resolution image synthesis." *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. 2021.
4. Clifford R Jack Jr et al. "The Alzheimer's disease neuroimaging initiative (ADNI): MRI methods". In: *Journal of Magnetic Resonance Imaging: An Official Journal of the International Society for Magnetic Resonance in Medicine* 27.4 (2008), pp. 685–691.

Realização



Apoio

