



Facultad de
Ingeniería de Sistemas

LAJC

LATIN-AMERICAN JOURNAL OF COMPUTING

Volume 11, ISSUE 1
January 2024
ISSN: 1390-9266
e-ISSN: 1390-9134

LJJC

Vol XI, Issue 1, January 2024



<https://www.epn.edu.ec>

MISIÓN

La Escuela Politécnica Nacional es una Universidad pública, laica y democrática que garantiza la libertad de pensamiento de todos sus integrantes, quienes están comprometidos con aportar de manera significativa al progreso del Ecuador. Formamos investigadores y profesionales en ingeniería, ciencias, ciencias administrativas y tecnología, capaces de contribuir al bienestar de la sociedad a través de la difusión del conocimiento científico que generamos en nuestros programas de grado, posgrado y proyectos de investigación. Contamos con una planta docente calificada, estudiantes capaces y personal de apoyo necesario para responder a las demandas de la sociedad ecuatoriana.

VISIÓN

En el 2024, la Escuela Politécnica Nacional es una de las mejores universidades de Latinoamérica con proyección internacional, reconocida como un actor activo y estratégico en el progreso del Ecuador. Forma profesionales emprendedores en carreras y programas académicos de calidad, capaces de aportar al desarrollo del país, así como promover y adaptarse al cambio y al desarrollo tecnológico global. Posiciona en la comunidad científica internacional a sus grupos de investigación y provee soluciones tecnológicas oportunas e innovadoras a los problemas de la sociedad.

La comunidad politécnica se destaca por su cultura de excelencia y dinamismo al servicio del país dentro de un ambiente de trabajo seguro, creativo y productivo, con infraestructura de primer orden.

ACCIÓN AFIRMATIVA

La Escuela Politécnica Nacional es una institución laica y democrática, que garantizar la libertad de pensamiento, expresión y culto de todos sus integrantes, sin discriminación alguna. Garantiza y promueve el reconocimiento y respeto de la autonomía universitaria, a través de la vigencia efectiva de la libertad de cátedra y de investigación y del régimen de cogobierno.



FACULTAD DE INGENIERÍA DE SISTEMAS

<https://fis.epn.edu.ec>

MISIÓN

La Facultad de Ingeniería de Sistemas es el referente de la Escuela Politécnica Nacional en el campo de conocimiento y aplicación de las Tecnologías de Información y Comunicaciones; actualiza en forma continua y pertinente la oferta académica en los niveles de pregrado y postgrado para lograr una formación de calidad, ética y solidaria; desarrolla proyectos de investigación, vinculación y proyección social en su área científica y tecnológica para solucionar problemas de transcendencia para la sociedad.

VISIÓN

La Facultad de Ingeniería de Sistemas está presente en posiciones relevantes de acreditación a nivel nacional e internacional y es referente de la Escuela Politécnica Nacional en el campo de las Tecnologías de la Información y Comunicaciones por su aporte de excelencia en las carreras de pregrado y postgrado que auspicia, la calidad y cantidad de proyectos de investigación, vinculación y proyección social que desarrolla y su aporte en la solución de problemas nacionales a través del uso intensivo y extensivo de la ciencia y la tecnología.

LAJC **LATIN-AMERICAN
JOURNAL OF
COMPUTING**

Vol XI, Issue 1, January 2024

ISSN: 1390-9266 e-ISSN: 1390-9134

Published by:
Escuela Politécnica Nacional
Facultad de Ingeniería de Sistemas

Quito – Ecuador

Indexed in



Associated institutions





Mailing Address

Escuela Politécnica Nacional,
Facultad de Ingeniería de Sistemas
Ladrón de Guevara E11-253, La Floresta
Quito-Ecuador, Apartado Postal: 17-01-2759

Web Address

<https://lajc.epn.edu.ec/>

E-mail

lajc@epn.edu.ec

Frequency

2 issues per year

Published by

Escuela Politécnica Nacional
Facultad de Ingeniería de Sistemas
Ecuador

Editor in Chief

Denys A. Flores. PhD.
Escuela Politécnica Nacional, Ecuador

Editorial Committee

Gabriela Suntaxi, PhD. (Chair)
Escuela Politécnica Nacional, Ecuador
gabriela.suntaxi@epn.edu.ec
Shahrazad Zargari, PhD.
Sheffield Hallam University, England
S.Zargari@shu.ac.uk
Matthew Bradbury, PhD.
University of Lancaster, England
m.s.bradbury@lancaster.ac.uk

Hagen Lauer, PhD.
Fraunhofer SIT, Germany
hagen.lauer@sit.fraunhofer.de
Diana Ramírez PhD (c).
Universidad Pompeu Fabra, España
diana.ramirez@upf.edu

Co-Editors

Carlos Iñiguez, Ph.D.
Escuela Politécnica Nacional, Ecuador
carlos.iniguez@epn.edu.ec
Iván Carrera, PhD.
Escuela Politécnica Nacional, Ecuador
ivan.carrera@epn.edu.ec

Edison Loza, Ph.D.
Escuela Politécnica Nacional, Ecuador
edison.loza@epn.edu.ec

Assistant Editors

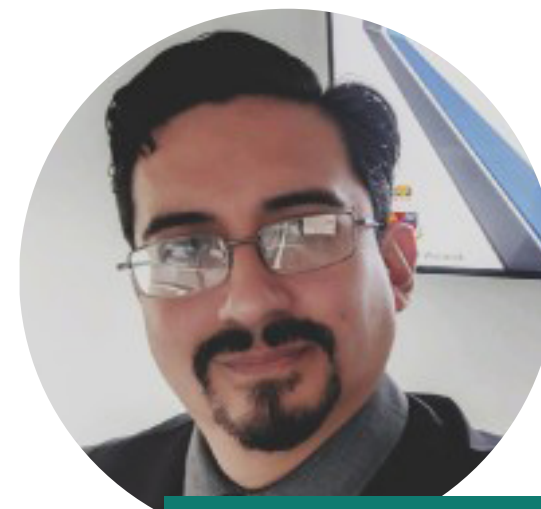
Ing. Damaris Tarapues
Technical Support
Escuela Politécnica Nacional, Ecuador
blanca.tarapues@epn.edu.ec

Ing. Gabriela Quiguango
Communications Manager
Escuela Politécnica Nacional, Ecuador
jenny.quiguango@epn.edu.ec

Proofreader

María Eufemia Torres, MSc.
Escuela Politécnica Nacional, Ecuador
maria.torres@epn.edu.ec

EDITORIAL



Denys Flores
PhD.

Editor LAJC
Escuela Politécnica Nacional,
Ecuador

Liderando la Innovación a través de las Aplicaciones de las Ciencias de la Computación

La constante evolución de las Ciencias de la Computación desafía a los investigadores a ampliar los límites de la innovación dentro de un panorama multidisciplinario. Desde la Editorial de la Revista Latin-American Journal of Computing, tenemos el agrado de presentar a nuestros lectores este número, el cual muestra investigaciones de vanguardia en diferentes aplicaciones de este campo.

El primer artículo explora la aplicación de los teoremas de Pappus-Guldin en el modelado de sólidos mediante interpolación spline. Aquí, los investigadores demuestran el potencial del análisis matemático para ofrecer soluciones informáticas más rentables para posiblemente optimizar el diseño de envases industriales. De manera similar, en el segundo artículo, se utiliza el modelado numérico para superar las limitaciones de los enfoques tradicionales en el análisis de la mecánica de fracturas elásticas lineales, comparando los resultados obtenidos utilizando plataformas comerciales y de código abierto.

Por el contrario, el trabajo presentado en el tercer artículo presenta esquemas esencialmente no oscilatorios para comprender el flujo de fluidos de dos fases en escenarios de extracción de petróleo. Los métodos numéricos se emplean con éxito para analizar los perfiles de mezcla de agua saturada y fluidos derivados del petróleo, lo que demuestra su importancia para comprender la dinámica de fluidos en materiales porosos. Así mismo, el cuarto artículo explora el uso de la optimización por enjambre de partículas para mejorar la eficiencia de un diseño de motor monofásico de reluctancia variable. Los autores demuestran que es posible minimizar las pérdidas de cobre mediante el análisis del método de elementos finitos.

El enfoque del quinto artículo aborda el panorama cambiante de la ciberseguridad. Los autores presentan una metodología para categorizar y actualizar ataques a servicios web, lo que contribuye a una mejor comprensión de las vulnerabilidades para prevenir estos ataques. Adicionalmente, el sexto artículo analiza la optimización de la asignación de recursos en Cloud Computing mediante la predicción del flujo de tráfico. Los investigadores emplean modelos de aprendizaje automático como ARIMA, Monte Carlo y XGBoost para dicho análisis predictivo.

Finalmente, los artículos séptimo y octavo cubren el diagnóstico médico y las necesidades educativas, respectivamente. En el primero, se presenta un método de diagnóstico temprano del Alzheimer mediante resonancia magnética y el algoritmo VGG16. Los autores justifican la eficacia del empleo de la IA para ayudar al diagnóstico de dicha enfermedad con una capacidad superior al 82 por ciento. En el último artículo, se utilizan técnicas de aprendizaje automático y minería de textos para explorar recursos educativos abiertos (REA) para identificar tópicos automáticamente, mejorando su descripción y categorización.

En conclusión, los artículos presentados en este número brindan una perspectiva única de las diferentes aplicaciones de las Ciencias de la Computación y la naturaleza dinámica de la investigación llevada a cabo en esta disciplina contemporánea. Gracias a los autores que contribuyeron al creciente cuerpo de conocimiento de este campo, deseándoles a ellos y a todos nuestros lectores un exitoso año 2024.

“Dejemos que la ciencia sea el vehículo para llevar nuestros sueños más allá de los límites de nuestra imaginación”.

Denys A. Flores

Editor en Jefe

Leading Innovation through the Applications of Computer Science

The constant evolution of Computer Science challenges researchers to push the boundaries of innovation within a multidisciplinary landscape. From the Editorial of the Latin-American Journal of Computing, we are pleased to present to our readership this number, which showcases cutting-edge research in different applications of this field.

The first article explores the application of Pappus-Guldin Theorems in solid modeling using spline interpolation. Here, researchers demonstrate the potential of mathematical analysis in order to deliver more cost-effective computing-based solutions to possibly optimize industrial packaging design. Similarly, in the second article, numerical modeling is used to overcome the limitations of traditional approaches for analyzing linear elastic fracture mechanics by comparing the results obtained using commercial and open-source platforms.

Conversely, the work featured in the third article presents essentially non-oscillatory schemes for understanding the flow of two-phase fluids in oil extraction scenarios. Numerical methods are successfully employed to analyze mixing profiles of saturated water and petroleum fluids, demonstrating their importance for understanding fluid dynamics in porous materials. Likewise, the fourth article explores the usage of particle swarm optimization for enhancing the efficiency of a single-phase variable reluctance motor design. The authors demonstrate that minimizing copper losses is possible through finite element method analysis.

Addressing the evolving landscape of cybersecurity is the focus of the fifth article. The authors introduce a methodology for categorizing and updating attacks on web services, which contributes to a better understanding of vulnerabilities for preventing web-based attacks. In addition, the sixth article discusses the optimization of resource allocation on Cloud Computing by predicting traffic flow. The researchers employ machine learning models like ARIMA, Monte Carlo, and XGBoost for such predictive analysis.

Finally, the seventh and eighth articles cover medical diagnosis and educational needs, respectively. In the former, an early-diagnosis method for Alzheimer's is featured using magnetic resonance imaging and the VGG16 Algorithm. The authors justify the effectiveness of employing AI to aid the diagnosis of such disease with a capacity exceeding 82 per cent. In the latter, machine learning and text mining techniques are used to explore open educational resources (OER) for automatically identifying topics, enhancing their description and categorization.

In conclusion, the articles brought to you in this number provides a unique perspective to the different applications of Computer Science, and the dynamic nature of the research carried out in this contemporary discipline. Thanks to the authors who contributed to the ever-growing body of knowledge in this field, wishing them, and all our readers, a successful year 2024.

“Let science be the vessel to carry our dreams beyond the limits of our imagination”

Denys A. Flores

Editor-in-Chief



We are most grateful to the following individuals for their time and commitment to review manuscripts for the Latin American Journal of Computing – LAJC

Reviewers

- Carlos Montenegro, MSc.**
Escuela Politécnica Nacional
- Daniela Buske, PhD.**
Universidade Federal de Pelotas
- David Nuñez, MSc.**
CNT EP
- Diego Almeida, PhD.**
Universidad Yachay Tech
- Elizaldo Domingues, PhD.**
Federal University of Rio Grande - FURG
- Emanuel Estrada, PhD.**
Federal University of Rio Grande - FURG
- Evelyn Mosquera, MSc.**
Escuela Politécnica Nacional
- Fran Lobato, PhD.**
Universidade Federal De Uberlândia
- Hernán Ordoñez, MSc.**
Escuela Politécnica Nacional
- Jorge Miño, MSc.**
Escuela Politécnica Nacional
- Jorge Zambrano, PhD.**
Universidad Politécnica de Valencia
- Julio dos Anjos, PhD.**
Universidade Federal do Rio Grande do Sul

- Leonardo Valdivieso, PhD.**
Escuela Politécnica Nacional
- Liércio Isoldi, PhD.**
Federal University of Rio Grande - FURG
- Marcelo Garcia, PhD.**
Universidad Técnica de Ambato
- Maritzol Tenemaza, PhD.**
Escuela Politécnica Nacional
- Monserate Intriago, PhD.**
Escuela Politécnica Nacional
- Olivia Barrón, PhD.**
Universidad de Monterrey
- Patricio Zambrano, PhD.**
Escuela Politécnica Nacional
- Raúl Córdova, MSc.**
Escuela Politécnica Nacional
- Régis Quadros, PhD.**
Universidade Federal de Pelotas
- Roberto Andrade, PhD.**
Escuela Politécnica Nacional
- Silvia Fajardo, PhD.**
Universidad de Colima
- Tania Calle, PhD.**
Escuela Politécnica Nacional
- Vera Ferreira, PhD.**
Universidade Federal do Pampa

TABLE OF CONTENTS

Pappus-Guldin theorems applied to the study of solid modeling with GeoGebra software

Fernanda Zastrow Tavares
Luverci Do Nascimento Ferreira

20

Numerical Modeling For Fracture Mechanics Problems Using The Open-source Fenics Platform

Caio César Prates Martins
Brenda Resende Lemos
Jean Pierre de Oliveira Bone
André Gustavo de Souza Galdino
Rodolfo Giacomim Mendes de Andrade

30

Essentially non-oscillatory schemes applied to Buckley-Leverett equation with diffusive term

Raphael de Oliveira Garcia
Graciele Paraguaia Silveira

42

Single Phase Variable Reluctance Motor Design Using Particle Swarm Optimization

Danyelle Schumanski
Dullian Carly de Oliveira Macedo
Juliana Almansa Malagoli
Cinthia Schimith Silva Coelho

56

The Attack Taxonomy Methodology Applied to Web Services

Marcelo I. P. Salas

66

An approach for optimizing resource allocation and usage in cloud computing systems by predicting traffic flow

Sello Prince Sekwatlakwatla
Vusumuzi Malele

80

Alzheimer's diagnosis system based on magnetic resonance imaging using the VGG16 algorithm

Werner Shtaniklao Ucañay Barreto
Marco Antonio Coral Ygnacio

90

Exploring Topics in Information Technology Open Educational Resources through the LDA Algorithm

René Ludeña
Verónica Segarra-Faggioni
Audrey Romero-Peláez
Juan Carlos Morocho-Yunga

106

Pappus-Guldin theorems applied to the study of solid modeling with GeoGebra software

ARTICLE HISTORY

Received 14 March 2023
Accepted 12 May 2023

Fernanda Zastrow Tavares
Universidade Federal do Rio Grande
(FURG)
Rio Grande, Brasil
fernandazastrow@gmail.com
ORCID: 0000-0001-8323-4221

Luverci do Nascimento Ferreira
Universidade Federal do Rio Grande
(FURG)
Rio Grande, Brasil
luverci@gmail.com
ORCID: 0000-0002-0662-9112

Pappus-Guldin theorems applied the study of solid modeling with GeoGebra software

Fernanda Zastrow Tavares
Instituto de Matemática, Estatística e Física (IMEF)
Universidade Federal do Rio Grande (FURG)
Rio Grande, Brasil
fernandazastrow@gmail.com
ORCID: 0000-0001-8323-4221

Luverci do Nascimento Ferreira
Instituto de Matemática, Estatística e Física (IMEF)
Universidade Federal do Rio Grande (FURG)
Rio Grande, Brasil
luverci@gmail.com
ORCID: 0000-0002-0662-9112

Abstract— In this work, we use Geogebra software to simulate the shape of objects (solids) in three dimensions from their photo and real dimensions using spline interpolation. With the reconstructed object, we analyze its volume and surface area using the Pappus-Guldin Theorems (PGT), the theorems that use mathematical analysis ideas to describe the volume and surface area by the sectional area and by the contour curve of the object. In the simulations, we tested the verification of the modeling for known solids (sphere and torus) and then analyzed some objects used in the industry, such as the packaging of products, pet bottles, yogurt containers, coffee powder packaging, aluminum soda cans, and the packaging of chocolate powder. We also analyzed some objects created by rotating bodies, such as the shape of a jar and an aluminum barrel, and also shapes found in nature, such as the shape of a pear and an egg. Modeling allows us to better understand the packaging used in the industry to minimize manufacturing costs and maximize its utility. Thus, we can modify these packages to obtain the best development of how these products are presented to the public, optimizing its format by analyzing its surface and its volume.

Keywords—Numerical simulation, GeoGebra, Spline, Pappus-Guldin theorems.

I. INTRODUCTION

Applications involving solid geometric objects that relate to the measurement of dimension or volume of a three-dimensional shape are classic examples that, according to [12], require the study of geometry of these bodies. The simple application of this work would be to calculate the area and volume of any three-dimensional object. This theoretical tool that gives us the volume and surface area of bodies, with the idea centered on axis rotation, are the Pappus-Guldin Theorems (PGT). Automatic fulfillment of the requirements for the use of electronic computers. This tool helps to understand how computational modeling is performed in certain solid models.

Based on the article [3] on rotational shapes and volume calculation, we use the numerical simulation code. The application will be a calculation of area and volume by computational modeling, starting from the analysis of an object from a photograph, which is inserted into the GeoGebra program according to [6] and using the spline command, can be identified the generation of contours of the object, building points on the contours of the image is called a polygonal line. Therefore, we will use the free software GeoGebra to model some objects and find the volume and surface area of these objects,

the axis can be formed by rotation around a surface or curve. Throughout the process, we have the minimum surfaces, and surfaces of revolution have a limit that will be discussed so that we can use the Pappus-Guldin theorems (PGT).

We will study these applications from the calculation of surface area and volume. For example, two methods can be compared: organic and industrial composting. In biological and engineering applications, we can use the PGT model to calculate the lateral surface area, and we can know the mass in terms of light and amount of water it produces on the surface in terms of the best fruit, flavor, and size. The industrial application of packaging can help to recycle and is less harmful to environment. We can minimize the product from the packaging in formats: for less heat exchange at room temperature; lower material costs. In the last part of the work, theoretical studies were carried out on the smallest surfaces of the packaging to minimize the design process of the designers of the formats of these packages.

II. METHODOLOGY

Numerical simulations of geometric solids using the mathematical theory of PGT will be applied. This can result in a study of minimizing the manufacturing costs of the products of companies through the packaging format of their products.

The software used for the calculations, Geogebra, is free software with mathematical functions developed by mathematician Markus Hohenwarter as part of his Ph.D. thesis [6] at the University of Salzburg. The program aims to develop suitable tools for teaching mathematics and applying mathematics in the fields of geometry and algebra.

In class, we use mathematical and physical concepts to follow theory and practice. However, nowadays they are not always together. There are often people who have the theory and do not know how to put it into practice or vice versa. Therefore, it shows the reality applied in daily practice

In this article, based on a new technology called Information and Communication Technology (ICT), we use this software to perform numerical simulations. GeoGebra in which we will model command objects with splines according to [4].

For theoretical support, we used the PGT to calculate the surface area and volume of solids. These theorems were proved by the Greek mathematician Pappus of Alexandria and the Swiss mathematician Paul Guldin, who used the concept of the center of rotation of solids to perform analyses and their measurements (see [9], [12], and [13]).

A. Pappus-Guldin theorems

There are two theorems referring to Pappus and Guldin:

First Pappus-Guldin theorem: Considering R as a plane figure, the solid Γ formed by rotation of R around an axis r has the volume of the solid is equal to area of R multiplied by the length of circle described by its center of gravity (see Fig.1).

The volume of the solid is thus given by the equation:

$$Vol(\Gamma) = 2\pi \text{dist}(r, G) A(R), \quad (3)$$

where G is the center of gravity of R , $2\pi \text{dist}(r, G)$ is the distance of the line r from the center of gravity, and $A(R)$ is the area of the region R .

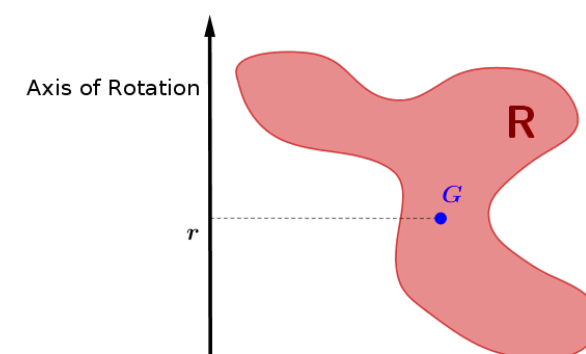


Fig.1. First Pappus-Guldin theorem

Second Pappus-Guldin theorem: Considering C as a plane curve, the surface Σ formed by rotation of C around an axis r has area obtained by multiplying the length of the curve C by length of the circumference passing through its center of gravity G (see Fig.2).

The surface area of Σ is given by the equation:

$$A(\Sigma) = 2\pi \text{dist}(r, G) l(C), \quad (4)$$

where G is the center of gravity of curve C , $2\pi \text{dist}(r, G)$ is the distance of the center of gravity of the lines r , and $l(C)$ is the length of curve C .

B. Spline

A spline is defined as a partitioned domain of a polynomial of degree n whose function value and its $n-1$ continuous first derivatives pass through the connection points. The abscissa of these connecting points are called knots, and these

piecewise polynomials are chosen to minimize the least mean square curvature.

According to [1], splines can be divided into two categories:

- Interpolation splines, which pass through all control points.
- Approximation splines, which run near all control points

Let $a = x_0 < x_1 < \dots < x_n = b$, be a subdivision of the interval (a, b) . A spline function of degree n with knots at the points x_i , $i=0, 1, \dots, m$ is a function S with the following properties according to [1]:

- In each subinterval (x_i, x_{i+1}) , $i=0, 1, \dots, m-1$, $S(x)$ is a polynomial of degree n .
- $S(x)$ and its first derivatives $(n-1)$ are continuous on the interval (a, b) .

According to [8], we can use spline functions to calculate tree volumes. The method used to calculate the volume of a tree would be to calculate the volume from the diameter and height of the trunk, using a cubic spline.

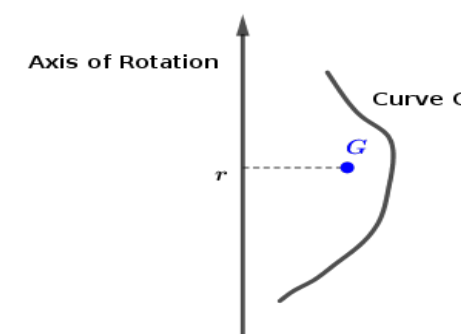


Fig.2. Second Pappus-Guldin theorem

III. NUMERICAL SIMULATION I

In this section, we will perform simulations using GeoGebra software and the Pappus-Guldin theorem (PGT) theory in the study of the volume and surface area of sphere and torus to verify the accuracy of the method.

A. Sphere

The 3D modeling of the sphere (see Fig.3) was performed in GeoGebra software according to [6]. We note that as the value of n increases, there is a built-in convergence for the sphere with radius $r=4$. Given a value and knowing that r is the fixed radius, we conclude that the volume of a sphere using the PGT (V_{PGT}) in the limit, approximates the volume value of the sphere using known results, i.e.,

$$\lim_{n \rightarrow \infty} V_{PGT} = V_{sphere} = \frac{4\pi r^3}{3} \quad (1)$$

where the locus of r is constant according to [5]. These values of absolute error and relative error decrease when only the value of n increases. We applied the same procedure by calculating a rough estimate of what was done for volume of the sphere with a different algorithm for the area of sphere using GeoGebra software from the creator of the application [6], which confirms the same idea.

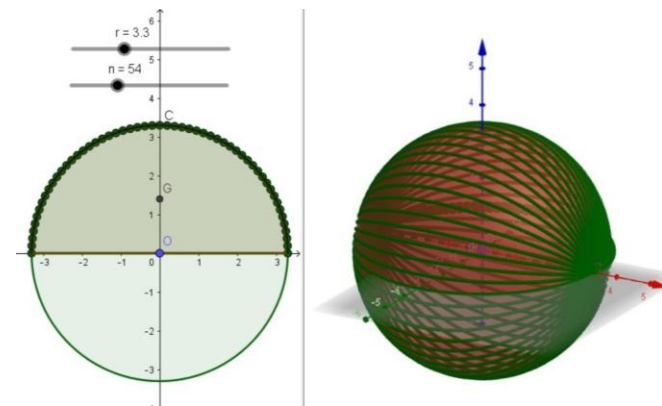


Fig. 3. Sphere simulation

Table I shows the algorithm of the PGT, which calculated the volume of the sphere in GeoGebra software. In the numerical simulation, we have a fixed radius of $r=4$, where we observe that the number of sides of the polygon inscribed in the circle depends on the value of n . The values of n are chosen randomly. Table I analyzes the value of the volume of the sphere, which is estimated to be the $V_{sphere} = 268.08$ u.v. Then we compare the value obtained over PGT volume with the modeling, which converges when it is verified that the value of n increases. We have below Table I:

TABLE I. SPHERE SIMULATION

Table 1: Sphere simulation			
n	V_{PGT}	Absolute error	Relative error
54	267.86	0.22	0.08206505521
77	267.97	0.11	0.0410325276
91	268	0.08	0.02984183826
111	268.03	0.05	0.01865114891

B. Torus

The 3D modeling of the torus (see Fig.4) was performed in GeoGebra software. The Cartesian plane coordinates x , y , and z in three dimensions represented by 3D is in the interval $[0, 2\pi]$, the torus has its rotational symmetry in the z -axis, R is the distance from the center of the tube to the center of the torus, and r is the radius of the tube.

The volume of the torus is given by:

$$V_{torus} = 2\pi^2 R r^2 = (\pi r^2)(2\pi R) \quad (2)$$

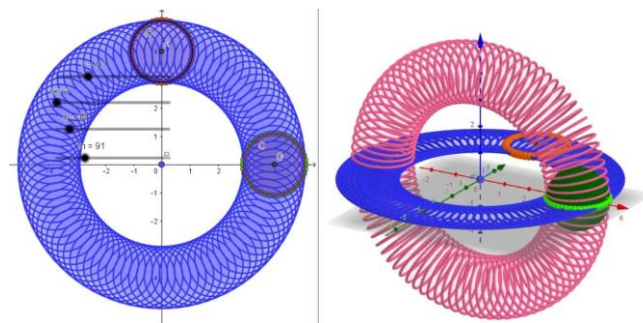


Fig. 4. Torus simulation

Table II shows the PGT algorithm performed in GeoGebra software for volume of a torus. In the numerical simulation, we have fixed radius of r and s , while the values of n and m can vary, i.e., the number of sides of the inscribed polygon formed inside the circle. The values of n and m are chosen randomly. Table II analyzes the value of the volume of a torus estimated at $V_{torus} = 3158.27341$. Then we compare the value of the PGT volume with the modeling, which converges when it finds that the values of n and m increase.

Below is Fig. 2 of a torus and table II:

TABLE II. TORUS SIMULATION

Table II: Torus simulation				
n	m	V_{PGT}	Absolute error	Relative error
39	37	3143.12	15.15	0.4796929965
91	44	3147.55	10.72	0.3394263315
107	195	3157.73	0.54	0.01709796819
242	227	3157.87	0.40	0.01266516162

IV. NUMERICAL SIMULATION II

In this section, we will perform simulations using Geogebra software and PGT theory to study the volume and surface area of some solids designed for industrial use, such as packaging and solids found in nature, such as eggs and pears.

A. Simulation in the packaging industry

It can also be shown that there are other objects used in industry, mainly in packaging, so we minimize the cost by optimizing the surface area and occupied volume with GeoGebra software to perform the modeling and analysis. Figures 5 and 6 show two minimal surfaces used as the basis for the production of some packaging, the catenoid, and the helicoid.

Numerical simulations are created by objects of geometric models with mathematical concepts using computer graphics. With its importance in reducing production time and costs, with a potential impact on job creation and income, as seen in [8]. Based on the use of software for manufacture of electronic components, we have used geometric modeling with mathematics according to [10], where the PGT can be applied to determine its volume and surface area in other products used in the industry.

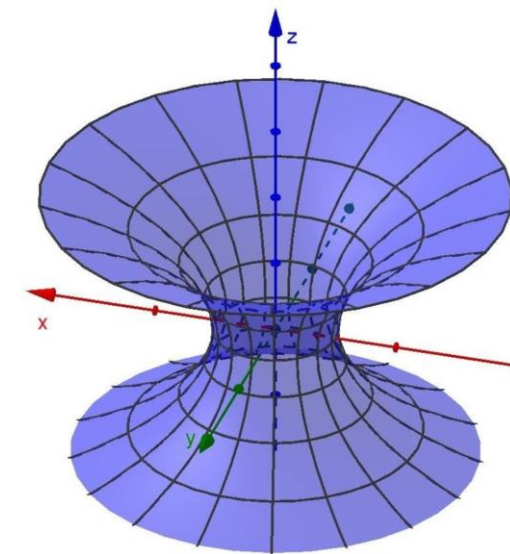


Fig.5. Catenoid

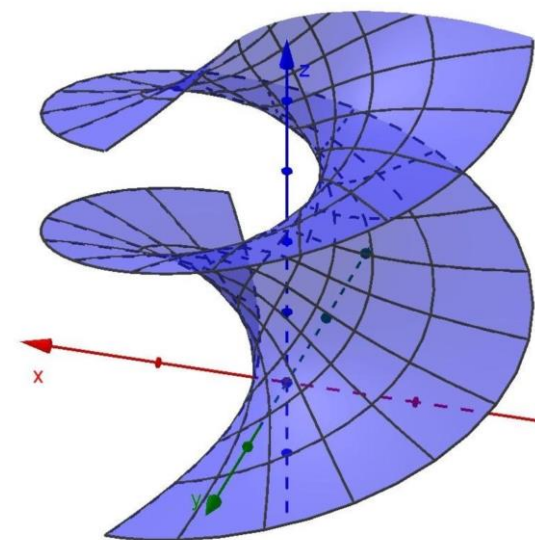


Fig. 6. Helicoid

Industry uses packaging to store its products, so it is important to study the format to be used. Reduce packaging costs by paying attention to factors such as format; reduce material costs. Among other things, analyzing these factors can lead to reducing environmental impact and promoting sustainable recycling. We can consider the use of minimal surface areas in industrial packaging. The importance of studying packaging forms, which we refer to in the work of [7], explores the feasibility and importance of new forms of PET bottles. In this case, when designing new packaging, we should primarily consider reducing the environmental impact of packaging damage. The examples of these packages by designers of PET bottle and yogurt figures are surfaces of revolution that have an axis of rotation, but they are not minimal

surfaces. To be a minimal surface, it would have to have infinite curves that fix two points on each edge.

B. Numerical models of packaging in the industry

Using real images of objects and using the resources of the Geogebra software, we were able to model some packages used in daily life.

Figures 7 to 13 present simulations considering objects modeled by using splines where the PGT was applied.

Fig. 7 has the modeling of the PET bottle with the number of steps $n=150$ and the value of the volume over $V_{PGT} = 25.16$ and the value over PGT of the area is $A_{PGT} = 30.2$.

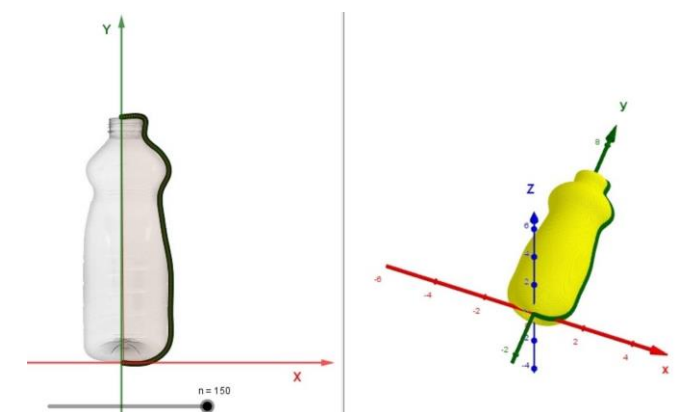


Fig.7. PET bottle

Fig. 8 has the modeling of a yogurt container with number of steps $n = 150$ and volume value over $V_{PGT} = 78.75$ and value over PGT of the area is $A_{PGT} = 62.98$.

Fig. 9 has the modeling of coffee powder packaging as a solid of revolution and a catenoid. The packaging is improved by the designer to attract consumer attention and increase sales through a naturally developed format.

From the visual point of view of packaging, the best would be the catenoid, as ours is cheaper to produce and the aluminum can has a strong construction and also for better conservation.

Modeling of coffee powder packaging with the number of steps $n = 150$ and volume value over $V_{PGT} = 48.59$ and value over PGT of the area is $A_{PGT} = 46.91$.

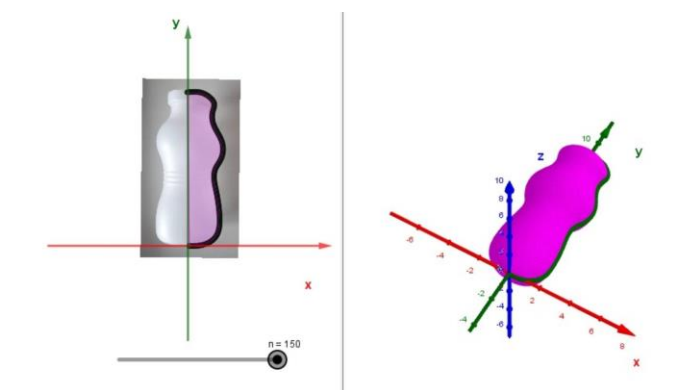


Fig. 8. Yogurt container

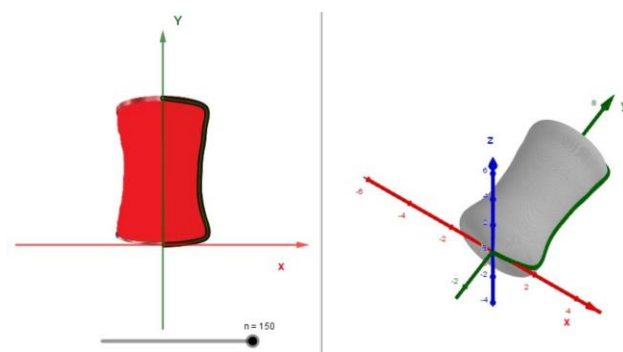


Fig. 9. Coffee powder packaging

Fig. 10 has the modeling of a jar with the number of steps $n = 150$ and the volume value over $V_{PGT} = 2.51$ and value over PGT of the area is $A_{PGT} = 16.19$.

In industry, the packaging of chocolate powder looks like a spiral, but it is not a spiral. This kind of filling cannot be modeled with PGT because we cannot get a generating curve around a fixed axis with the spline, in short, it is not a surface of revolution.

Some packages are modified to achieve the best development of a product to attract the public and sell more products through advertising [2].

In the studied example, the chocolate milk can is not optimized, and the production cost of can must have increased.

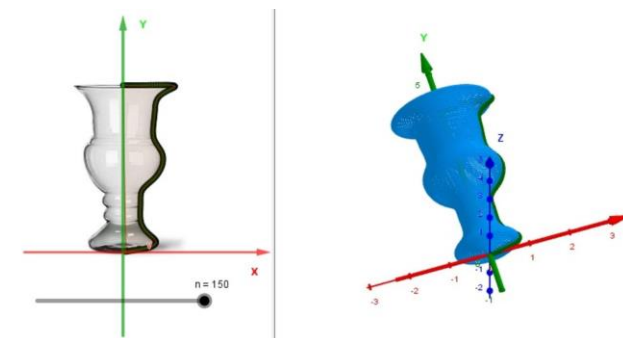


Fig. 10. Jar

Oblique cylindrical shapes are mainly used in the automotive industry, such as drills, screws, rotors, and helical gears, because the material is more evenly distributed and more resistant, which increases the quality of the product.

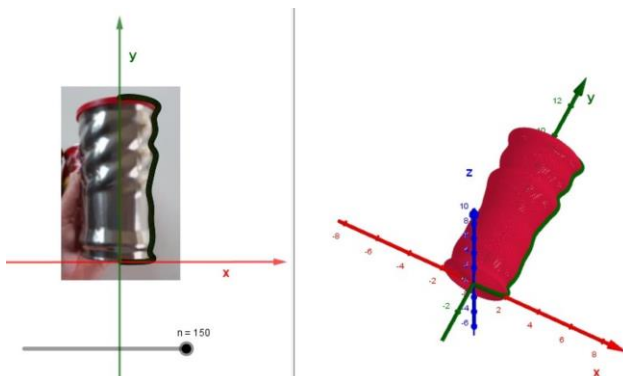


Fig. 11. Packaging of chocolate powder

We can model the packaging of chocolate powder if we assume that it is obtained by a generation curve. To do this, we can use the analysis of the model through PGT, as shown in Fig. 11.

Thus, we obtain a modeling with a number of the steps $n = 150$ and the volume value over $V_{PGT} = 102.58$ and the PGT value of the area is $A_{PGT} = 66.81$.

Fig. 12 has the modeling of an aluminum soda can with the number of steps $n = 150$ and the volume value over $V_{PGT} = 196.4$ and the PGT value for the area is $A_{PGT} = 111.86$.

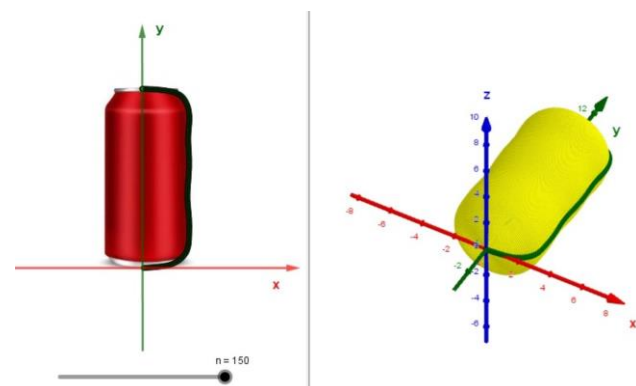


Fig. 12. Aluminum soda can

Fig. 13 shows the modeling of an aluminum barrel with a number of steps $n = 150$ and the volume value over $V_{PGT} = 95.65$ and the PGT value for the area is $A_{PGT} = 66.69$.

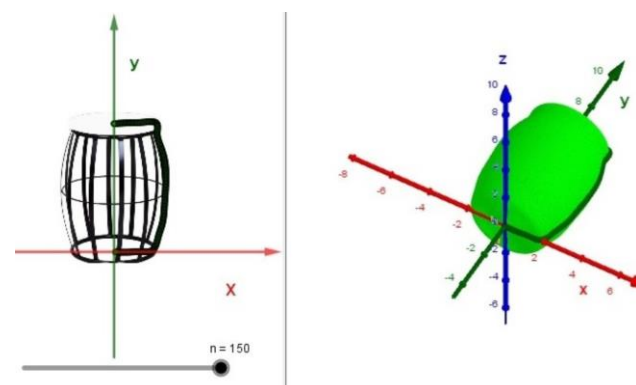


Fig. 13. Aluminum barrel

C. Numerical models of packaging in nature

Figures 14 and 15 show the format of some packages found in nature, such as pear and egg.

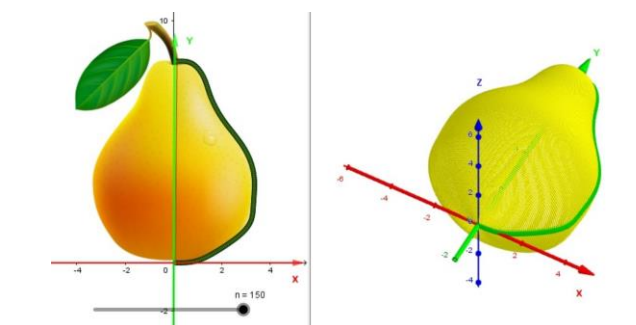


Fig. 14. Pear shape

REFERENCES

- [1] Anton, H., Rorres, C., “Álgebra Linear com aplicações”, Bookman, Porto Alegre, vol 8, p.16, 2001.
- [2] Benainous, Edileuza Coelho da Silva., “Nescau: uma análise da estratégia de Reposicionamento do produto via embalagens”.
- [3] Dantas, S.C., Mathias, C.V., “Formas de revolução e cálculo de volume”, Ciência e Natura, vol 39, n. 1, p. 142-155, 2017.
- [4] Andrade, M.S., Simulação e modelagem computacional com o software Modellus: aplicações práticas para o ensino de física”, Editora Livraria da Física, 2016.
- [5] Carmo, M.P., “Geometria diferencial de curvas e superfícies”, SBM-Sociedade Brasileira de Matemática, 2005.
- [6] Hohenwarter, M., (2002). “GeoGebra”, Available in: <http://www.geogebra.org/cms/en>. Accessed on June 25, 2022.
- [7] Keller, I., Vicente, F., dos Santos, R., “Um Novo Formato de Garrafa Pet”, XI Simpósio de Excelência em Gestão e Tecnologia, AEDB, 2014.
- [8] Kirchner, F.F., Filho, A.F., Scolforo, J.R.S., Machado, S.A., Mitishita, E.A., “O uso de funções spline no cálculo de volumes de árvores”, Floresta, vol 19, n. 1, UFPR, 1989.
- [9] Kurokawa, C.Y., “Áreas e volumes de Eudoxo e Arquimedes e Cavalieri e o Cálculo Diferencial e Integral”, UNICAMP, 2015.
- [10] Nunes, É.O., “Modelagem Geométrica e Sweeping”, Uff. Available in: <http://www.ic.uff.br/~aconci/sweeping.html>. Accessed on June 25, 2022.
- [11] Pinho, M.S., “Computação Gráfica Modelagem de Sólidos”, PUCRS. Available in: <https://www.inf.pucrs.br/~pinho/CG/Aulas/Modelagem/Modelagem3D.htm>. Accessed on July 25, 2022.
- [12] Rooney, A., “A História da Matemática: Desde a criação das pirâmides até a exploração do infinito”, São Paulo: M. Books do Brasil Editora LTDA, 2012.
- [13] Roque, T., “História da matemática”, Editora Schwarcz-Companhia das Letras, 2012.
- [14] Trevisan, E.P., Trevisan, A.C.R., “Teorema de Pappus-Guldin com auxílio dos softwares maxima e winplot: atividade a uma turma de cálculo II”, III Simpósio Internacional de Pesquisa em Educação Matemática – III SIPEMAT, UFC, 2012.
- [15] Volny, Petr., Volna, Jana., Smetanová, Dana., “Visualization of Numerical Data in Geogebra”, Littera Scripta, n.2, 2012

V. CONCLUSIONS

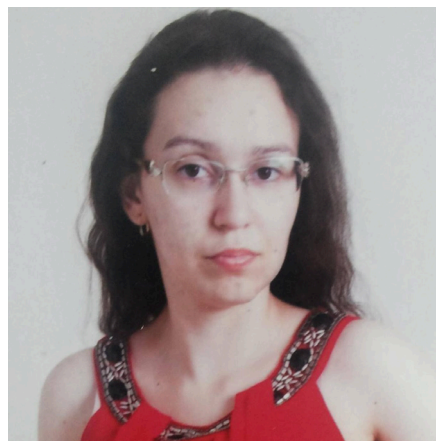
The modeling was done in free software GeoGebra. The formats of some parcels were constructed by rotation curves on the symmetry axis in both two and three dimensions. These Pappus-Guldin Theorems (PGT) allow us to measure and analyze the numerical value of these models, such as volume and surface area. The advantage of the PGT method is that it allows to analyze of numerical values of objects easily and directly with Geogebra, the disadvantage of the method is that it can be used only for solids obtained by rotation. The applied program provides a simple language as long as the figure has an axis of symmetry according to [14]. Since the work is extensive, we note that we combine theory and practice to provide practical results that can be applied in industries, such as the packaging industry.

In future work, costs can be minimized, especially in the production of packaging, by conducting a study on the use of minimum surface area to determine the shapes of some packaging, as well as a method to improve the use of packaging is to reduce the material used in the construction of this packaging and shaping to be able to reduce overall costs and bring better transportation, storage, distribution, and consumption conditions for the sale and consumption of goods. As a result, logistics would spend less time to reduce costs, a strategy to better serve customers, vehicles that consume less fuel, are less harmful to the environment, and have potential in their use of recycling in packaging.

ACKNOWLEDGEMENTS

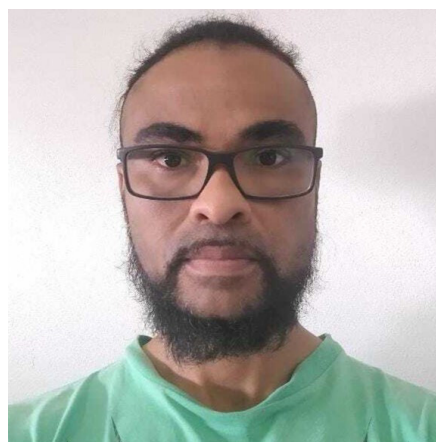
This work is the result of research carried out at IMEF - FURG (Institute of Mathematics, Statistics, and Physics of the Federal University of Rio Grande).

AUTHORS



Fernanda Tavares

Fernanda Zastrow Tavares, Brazilian, was born on April 9, 1992, in the city of Rio Grande-RS and is currently a student of Applied Mathematics at the Federal University of Rio Grande - FURG. She entered academic life in the Mathematics program at the Federal University of Rio Grande - FURG in 2010, graduated in 2016 and entered the Applied Mathematics program in 2017, with the aim of deepening her knowledge in the field of applied mathematics. Her area of interest is algebra, algorithm analysis, numerical simulations, mathematical modeling, and computer modeling in industry. In 2020, she began her studies in computer modeling, analyzing the Pappus-Guldin theorem and its applications to surfaces of revolution using Geogebra software and developing a monograph in applied mathematics. In October 2022, she presented the result of her research at the XXV National Meeting of Computational Modeling.



Luverci Ferreira

Luverci do Nascimento Ferreira, Brazilian, born in Brasília-DF, is currently an Adjunct Professor at the Institute of Mathematics, Statistics and Physics of the Federal University of Rio Grande - IMEF-FURG, since 2008. He has a degree in Mathematics from the University of Brasília - UnB (2003) and Master in Mathematics from the University of Brasília - UnB (2006) in the area of Ordinary Differential Equations. PhD student in Computational Modeling at the Federal University of Rio Grande since 2019, with a research area related to the study of algorithms to solve fractional differential equations using the statistical method known as the Monte Carlo Method. He has experience in teaching undergraduate courses in Mathematics, Applied Mathematics, Mechanical Engineering and Business Mechanical Engineering and in specialization courses for Mathematics Teachers. Research and interest in the areas of Mathematical Analysis, Mathematical Epidemiology, Biomathematics, Dynamic Geometry, Complex Variables, Ordinary and Partial Differential Equations, Fractional Calculus, Fractional Differential Equations, Stochastic Processes and Computational Modeling.

Numerical Modeling For Fracture Mechanics Problems Using The Open-source Fenics Platform

ARTICLE HISTORY

Received 15 March 2023
Accepted 12 May 2023

Caio César Prates Martins

Federal Institute of Espírito Santo
Vitória - Brazil
caiocp.martins@gmail.com
ORCID: 0009-0009-0609-200X

Brenda Resende Lemos

Federal Institute of Espírito Sant
Vitória - Brazil
brendarlemos2020@gmail.com

Jean Pierre de Oliveira Bone

Federal Institute of Espírito Santo
Guarapari - Brazil
jean.bone@ifes.edu.br

André Gustavo de Souza Galdino

Federal Institute of Espírito Sant
Vitória - Brazil
andre.galdino@ifes.edu.br
ORCID: 0000-0002-5990-0287

Rodolfo Giacomim Mendes de Andrade

Federal Institute of Espírito Santo
Vitória - Brazil
rodolfo.andrade@ifes.edu.br
ORCID: 0000-0003-3956-5941

Numerical Modeling For Fracture Mechanics Problems Using The Open-source Fenics Platform

Caio César Prates Martins
Federal Institute of Espírito Santo
Vitória - Brazil
caiocp.martins@gmail.com
ORCID: 0009-0009-0609-200X

Brenda Resende Lemos
Federal Institute of Espírito Santo
Vitória - Brazil
brendarlemos2020@gmail.com

Jean Pierre de Oliveira Bone
Federal Institute of Espírito Santo
Guarapari - Brazil
jean.bone@ifes.edu.br

André Gustavo de Souza Galdino
Federal Institute of Espírito Santo
Vitória - Brazil
andre.galdino@ifes.edu.br
ORCID: 0000-0002-5990-0287

Rodolfo Giacomim Mendes de Andrade
Federal Institute of Espírito Santo
Vitória - Brazil
rodolfo.andrade@ifes.edu.br
ORCID: 0000-0003-3956-5941

Abstract— Fracture mechanics is the mechanical approach to fracture processes, which emerged due to limitations in applying traditional concepts of Mechanics of Materials to predict the behavior of cracked materials. Analytical problem solutions with this approach may be unattainable, so it is necessary to use numerical modeling, such as the finite element method. However, the use of more advanced software that solves engineering problems numerically is limited by its high cost. FEniCS is an open-source computational platform that solves partial differential equations by the finite element method. Thus, from a tutorial for this computational platform, this work proposes to reproduce a classic problem of linear elastic fracture mechanics, based on the validation of a comparison of a linear elastic problem with the commercial software ANSYS®. With the help of the provided tutorial, a code was built to model a three-point bending test. Implemented with the aid of Gmsh and Paraview, it was possible to obtain satisfactory results and to show that FEniCS is a powerful and accessible tool for solving fracture mechanics problems.

Keywords— Fracture Mechanics, Numerical Modeling, FEniCS, Finite Elements

I. INTRODUCTION

The Finite Element Method (FEM) is an approach to solving partial differential equations using numerical techniques in which a continuous domain is discretized into finite elements called mesh. With the advancement of technology and, consequently, computational power, more advanced engineering problems have become simpler to solve. This is due to the fact that analytical solutions can be complex and even unreachable, and the error achieved in numerical solutions is considerably acceptable [1].

In general, the use of advanced FEM-based software is restricted to companies and some teaching institutions, as it is the case with ANSYS® [2]. On the other hand, open-source programs, such as the FEniCS Software [3], which is free, are used more widely due to their availability and ease of access. However, they tend not to have a graphical interface, unlike commercial software, which makes the use of pre- and post-processors essential for visualizing the solution. Even for a simple computational solution approach, the development of the finite element method can be complex, as it requires knowledge about tensor and variational calculus in some cases. The FEniCS [3] software, for example, uses these

principles and, based on programming knowledge and the library itself, it is possible to solve partial differential equations using a variational approach. Therefore, this work proposes to present a numerical modeling of a linear elastic problem and two of fracture mechanics, presenting a preliminary comparison with ANSYS® for the linear elastic problem.

A. Fracture Mechanics

Fracture mechanics is the mechanical approach to fracture processes that emerged due to limitations applying traditional concepts of Mechanics of Materials to predict the behavior of materials in the presence of cracks. It was developed and founded after the 2nd World War [4] and is widely used in structural contexts in the areas of civil, mechanical, and metallurgical engineering [5-7]. For materials with brittle behavior, the linear elastic fracture mechanics (LEFM) approach is used, while for materials with ductile behavior, the elastoplastic fracture mechanics (EPFM) is used [4,8].

For example, in industry, a component may have such a high cost that, depending on the conditions and its integrity, it is more viable to have knowledge about fracture mechanics and continue using it with cracks to perform an exchange, which results in a complete pause of an operation. Alan Arnold Griffith studied the behavior of an elliptical hole when external stress is applied and established a thermodynamic model for crack propagation [8]. Griffith concluded that the strength of a material is not only linked to chemical bonding parameters but also to the existing defects. Therefore, it was realized that defects in the material are factors that intensify the applied stress, making it susceptible to exceeding the yield strength of the material and causing a rupture, which is the basis of fracture mechanics [4]. With this, a good characterization of the material must also have experimental parameters of the LEFM, such as the fracture toughness (K_{IC}) and the critical energy release rate (G_c), for example. Fracture toughness is independent of size, geometry, and loading levels for a material with a given microstructure and is the main obtained experimentally properties related to fracture mechanics [8].

II. METODOLOGY

A. Numerical Modeling

The main concept addressed in the FEM is the discretization of a continuous domain into finite geometric elements, in addition to the use of polynomial interpolation to determine the results in the region inside the elements [1]. These elements form a mesh, and each node has a displacement u and a stress σ , which are represented in a linear system and determined through the variational calculus. The stress and strain of the solid are expressed by tensors, which are, by definition, mathematical entities that produce a linear transformation in vectors, transforming them into different vectors [9]. Tensors are represented in Equations (1) and (2), where u , v , and w are the horizontal, vertical, and transverse components of infinitesimal displacement. The strain tensor can also be defined as $\epsilon = \text{sym} \nabla u$, that is, the symmetric gradient of u . Both mathematical entities σ and ϵ have Cartesian x , y , and z components.

$$\sigma = \sigma^T = \begin{bmatrix} \sigma_{xx} & \sigma_{xy} & \sigma_{xz} \\ \sigma_{yx} & \sigma_{yy} & \sigma_{yz} \\ \sigma_{zx} & \sigma_{zy} & \sigma_{zz} \end{bmatrix} \quad (1)$$

$$\epsilon = \begin{bmatrix} \frac{\partial u}{\partial x} & \frac{1}{2} \left(\frac{\partial u}{\partial y} + \frac{\partial v}{\partial x} \right) & \frac{1}{2} \left(\frac{\partial u}{\partial z} + \frac{\partial w}{\partial x} \right) \\ \frac{1}{2} \left(\frac{\partial u}{\partial y} + \frac{\partial v}{\partial x} \right) & \frac{\partial v}{\partial y} & \frac{1}{2} \left(\frac{\partial v}{\partial z} + \frac{\partial w}{\partial y} \right) \\ \frac{1}{2} \left(\frac{\partial u}{\partial z} + \frac{\partial w}{\partial x} \right) & \frac{1}{2} \left(\frac{\partial v}{\partial z} + \frac{\partial w}{\partial y} \right) & \frac{\partial w}{\partial z} \end{bmatrix} \quad (2)$$

Mesh refinement generates more accurate results that are closer to the analytical ones, but there is a computational limit to be respected, which is analyzed through a convergence test, where the best results are sought with the minimum possible elements [9]. In the formulation involving the FEM applied to fracture mechanics, some parameters must be provided to the program, such as Young's modulus (E), Poisson's coefficient (ν), and the critical energy release rate (G_c). After providing the input data, using concepts of variational calculus, it is possible to obtain the results, which are observed through a post-processing software.

Simulation allows engineers to use basic principles of modeling, physics, mathematics, and computer science to evaluate design performance in different scenarios. Thus, for the development of Engineering, it is important to analyze solutions via software to ensure that the result obtained is adequate and that it meets the functional needs of a project [10].

1. Numerical modeling for FEniCS software

a. Linear elasticity

Numerical modeling is used for a linear elasticity problem in a plane strain state [9] shown in Fig. 1, based on the FEniCS library [11]. The problem consists of a three-dimensional plate 200 mm long, 500 mm high and with thick $e = 10$ mm subjected to a load. Acting field forces are disregarded and a plane strain state is defined. For the Young's modulus of the material, 200 GPa was adopted and, for Poisson's coefficient, 0.3. The problem has the following governing equations for a Ω domain.

$$-\nabla \cdot \sigma = T, \forall y = 0 \text{ in } \partial\Omega \quad (1)$$

$$\sigma = \lambda \cdot \text{tr}(\epsilon)I + 2\mu\epsilon \quad (2)$$

$$\epsilon = \frac{1}{2}(\nabla u + (\nabla u)^T) \quad (3)$$

where T is the applied stress, represented by the ratio between the uniformly distributed load at the base of the bar and the thickness, I is the three-dimensional identity matrix, and μ and λ are the Lamé constants, which depend on the Young's modulus and the Poisson's coefficient of the material. Considering the principle of virtual work, one must find values of u that satisfy the weak formulation [9].

$$\int_{\Omega} \sigma(u) : \epsilon(u) dx = \int_{\Omega} T \cdot p ds, \forall u, p \in V \quad (4)$$

where u and p are the trial and test functions, respectively, and V is the vector field containing them. In this example, second-degree Lagrange polynomials are defined for the interpolation between nodes. The vertical face is fixed, and the load F is uniformly applied in the negative y -direction. The mesh was built using a function from the FEniCS library, containing 48000 tetrahedral elements with 5 mm on each side.

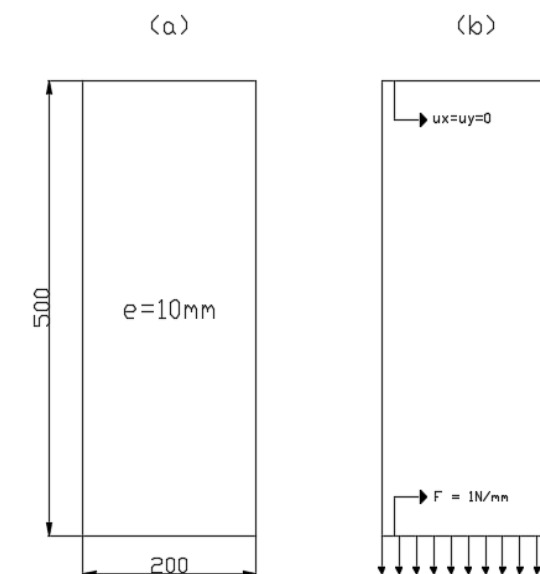


Fig. 1. Representation of geometry (a) and boundary conditions (b)

b. Fracture Mechanics

Numerical modeling for the fracture mechanics problem was proposed by [12], using models by [13], with contributions by [14]. It is considered an elasto-static body with a discontinuity, which occupies a domain $\Omega \subset \mathbb{R}^2$. The Dirichlet and Neumann boundary conditions [9] are imposed by Γ_D and Γ_C . In the case of a discrete fracture mechanism, the crack is represented by a discontinuous surface Γ_C . The variable that models crack propagation is $\phi \in [0,1]$. When it assumes a null value, the material is intact, and when it assumes a unitary value, there is a complete fracture. The crack size is controlled by a variable ℓ , a length scale parameter inherent to the model and which depends on the developed mesh refinement [12-15], called characteristic length. The approximate crack surface energy is defined as:

$$\int_{\Omega} G_c d\Gamma_c \approx \int_{\Omega} G_c \left(\frac{1}{2\ell} \phi^2 + \frac{\ell}{2} |\nabla \phi|^2 \right) d\Omega \quad (5)$$

Adding the Bulk energy to Eq. (5) the total potential energy of the solid (Ψ) is obtained as:

$$\Psi = \int_{\Omega} \left((1 - \phi^2) \psi(\varepsilon) + G_c \left(\frac{1}{2\ell} \phi^2 + \frac{\ell}{2} |\nabla \phi|^2 \right) \right) d\Omega \quad (6)$$

where $\psi(\varepsilon)$ is the strain energy density of the solid, in terms of the Lamé parameters and the strain tensor, represented in Eq. (2), whose mathematical expression is:

$$\psi(\varepsilon) = \frac{1}{2} \lambda (\text{tr}(\varepsilon))^2 + \mu \text{tr}(\varepsilon^2) \quad (7)$$

Applying Gauss's theorem in Eq. (7) the following field equations are obtained, with arbitrary values for the kinematic variables δu and $\delta \phi$.

$$\begin{aligned} (1 - \phi^2) \nabla \cdot \sigma &= 0, \text{ in } \Omega \\ G_c \left(\frac{1}{\ell} - \ell \Delta \phi \right) - 2(1 - \phi) \psi(\varepsilon) &= 0 \end{aligned} \quad (8)$$

The natural boundary conditions for a traction T are:

$$\begin{aligned} (1 - \phi^2) \sigma \cdot n &= T, \text{ on } \Gamma \\ \nabla \phi \cdot n &= 0, \text{ on } \Gamma \end{aligned} \quad (9)$$

where n is the normal vector to the surface Γ . With this, the constitutive equations and the boundary conditions are given. The procedure now consists of implementing the finite element method. The main objective is the resolution of the system of equations (8) with the boundary conditions expressed by Eqs. (9.1) and (9.2). However, it is necessary to use the finite element method, discretizing the continuous domain. Equation (8.2) is modified to:

$$G_c \left(\frac{1}{\ell} - \ell \Delta \phi \right) - 2(1 - \phi) H^+(\varepsilon) = 0 \quad (10)$$

where H^+ is called the variable storage (or history) field, which changes with time, expressed mathematically as:

$$H^+(\varepsilon) = \max \psi^+(\varepsilon(t)) \quad (11)$$

and ψ^+ is the variable strain energy density of the solid, defined as:

$$\psi^+(\varepsilon) = \frac{1}{4} K (\varepsilon + |\varepsilon|)^2 + \mu (\varepsilon^{\text{dev}} : \varepsilon^{\text{dev}}) \quad (12)$$

where K is the Bulk modulus, which can be expressed in terms of the Young's modulus and the Poisson's coefficient [9]. Finite element modeling uses a weak, or variational, formulation that uses dimensional trial ($\mathcal{U}, \mathcal{P}$) and test ($\mathcal{V}, \mathcal{Q}$) spaces, which contain the trial (u, ϕ) and test (p, q) functions, respectively. A discrete space ($\mathcal{W}$) is also defined around the mesh that contains the phase field variable (ϕ) and the displacement field (u). All spaces have a dimension d .

$$\begin{aligned} (\mathcal{U}, \mathcal{V}) &= \{(u, p) \in [C^0(\Omega)]^d : (u, p) \in [\mathcal{W}(\Omega)]^d \\ &\subseteq [H^1(\Omega)]^d \\ (\mathcal{P}, \mathcal{Q}) &= \{(\phi, q) \in [C^0(\Omega)]^d : (\phi, q) \in [\mathcal{W}(\Omega)]^d \\ &\subseteq [H^1(\Omega)]^d \end{aligned} \quad (13)$$

In the reformulation of the system of constitutive equations, applying the Bubnov-Galerkin procedure, remote tractions and field forces are disregarded, making it:

$$\begin{aligned} \int_{\Omega} [(1 - \phi)^2 \sigma(u) : \varepsilon(p)] d\Omega &= 0 \\ \int_{\Omega} [\nabla q \cdot \nabla \phi G_c \ell + q \left(\frac{G_c}{\ell} + 2H^+ \right) \phi - 2H^+ q] d\Omega &= 0 \end{aligned} \quad (14)$$

2. Numerical modeling for ANSYS ® software

In order to validate the results obtained by FEniCS for the linear elasticity problem (Section II.1.a), a numerical model was implemented in the Ansys ® software, version 2022. The material used in the modeling has the same properties as in

Section II.1.a. As it is a three-dimensional model, a load of 0.1N/mm² (Pressure type) was applied to the lower face of an xz plane of the structural element in the vertical direction with a downward direction (-y). Opposite the load application plane, all nodes were restricted to translation, which represented a crimp. The mesh illustrated in Fig. 2 records the geometry containing 45125 nodes and 8000 cubic elements.

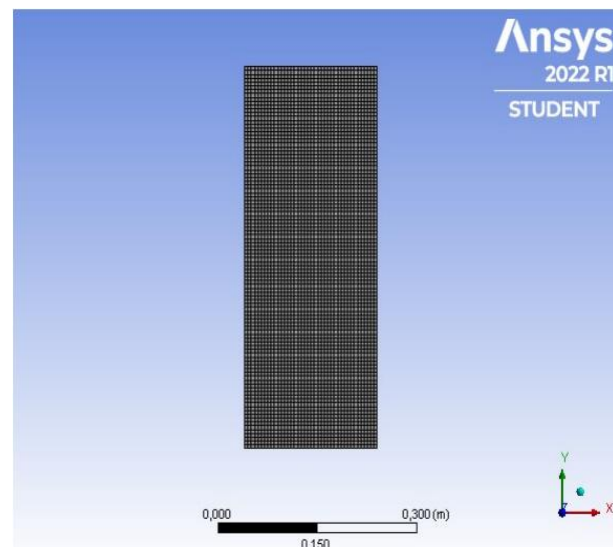


Fig. 2. Mesh result obtained in ANSYS ® for linear elasticity problem

III. RESULTS AND DISCUSSION

A. FEniCS software validation

With the help of ANSYS ®, a kind of FEniCS validation was carried out, solving the same linear elasticity problem and comparing the results. Fig. 3 shows the bar displacements obtained by ANSYS ®, and Fig. 4 exposes those obtained by FEniCS, both in the y direction. Note that there is a qualitative similarity regarding the vector field represented by the scale. The maximum supported stresses are found on the crimped face of the bar, opposite to the force application face, and, for both cases, a value of 0.1MPa was obtained. The maximum deformation obtained analytically is -2.500·10⁻⁴ mm. The comparison between FEniCS and ANSYS ® for the linear elastic problem resulted in errors of less than 0.6%, as shown in Table I. Therefore, free open-source software can be operated safely.

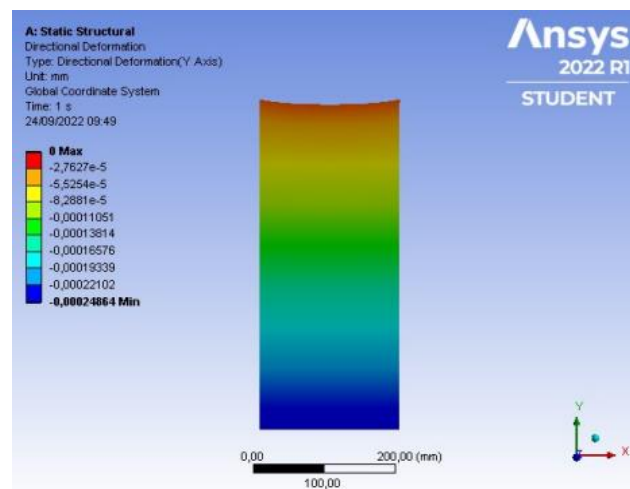


Fig. 3. Result of the displacement field obtained in ANSYS ® for linear elasticity problem

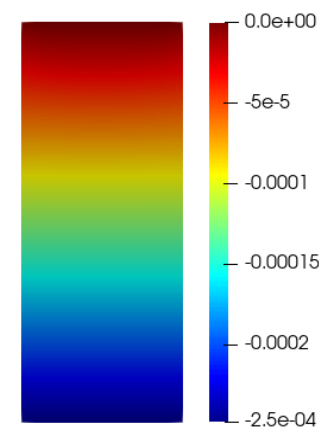


Fig. 4. Result of the displacement field obtained in FEniCS for linear elasticity problem

TABLE I. COMPARISON BETWEEN ANSYS AND FENICS MAXIMUM DEFORMATION

Ansys vs FEniCS comparison		
Method	Value (mm)	Relative Error ^a (%)
Ansys ®	-2.486·10 ⁻⁴	0.560
FEniCS	-2.487·10 ⁻⁴	0.520
Analytical	-2.500·10 ⁻⁴	-

^a. Relative to the analytical value

B. Application of the phase field method using FEniCS

1. Tensile test with simple pre-crack

Elaborated by [15], the problem to be solved consists of a representation of a fracture mechanics test of a plate subjected to uniaxial tensile stress that has a pre-crack to simulate a pure fracture in Mode I, as illustrated in Fig. 5. In order to reduce the computational time, small geometric proportions were considered, being $L=0.5$ mm. The code structure of this problem, implemented for this work, followed the tutorial developed by [12].

The mesh was built using the Gmsh preprocessor [16] and has 30546 triangular elements. The material has a modulus of elasticity $E = 210\text{GPa}$, a Poisson coefficient $\nu = 0.3$, and a critical energy release rate $G_c = 2.7\text{MPa mm}$. Thus, the Lamé parameters $\lambda = 121153.8\text{MPa}$ and $\mu = 80769.2\text{MPa}$ were obtained. A value of 0.011mm was also used for the characteristic length ℓ .

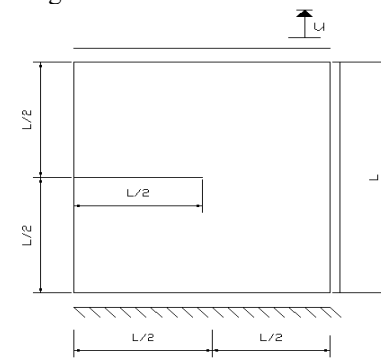


Fig. 5. Tensile stress problem in a pre-cracked plate under uniaxial force, adapted from [12,13]

The base of the mesh ($y=0$) is fixed, and a remote offset of 0.007mm is used as the first iteration. To help with the code, the value of the phase field variable ϕ for every pre-crack was defined as 1. During the execution of the code, the necessary number of iterations for convergence of the solution during a given step is provided. The main results of the analysis are shown in Fig. 6 and Fig. 7. Visualization of crack propagation is easily observed using the post-processor software Paraview [17]. Code execution stops at a value determined as a maximum ($t=1.0$), at which there has already been catastrophic failure of the material. A change was made regarding the test loading rate, for reasons of computational power. The red region represents the complete failure of the material, and the blue shows the initial state (intact).

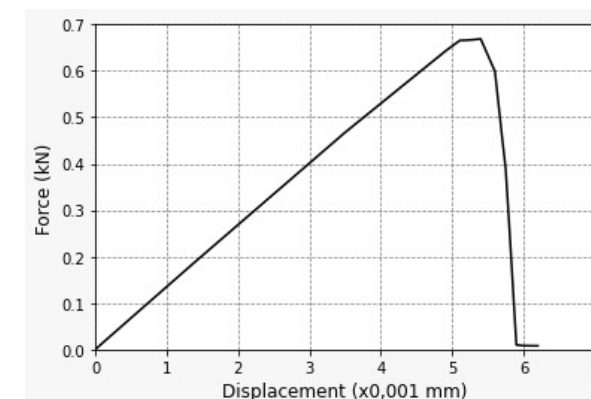


Fig. 6. Force-displacement curve for traction problem (a)

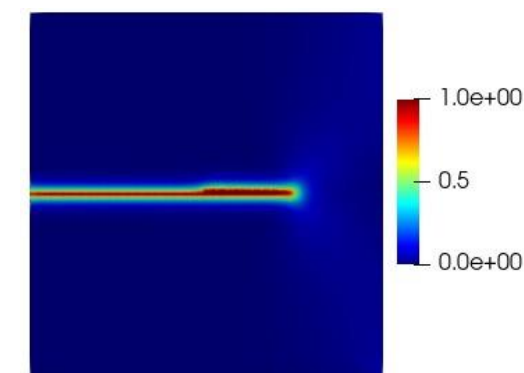


Fig. 7. Displacement crack propagation $u=5,7 \times 10^{-3}$ mm

The results obtained are satisfactory and close to the literature, as shown in Table II [12,13,15,18], although changes have been made to the code and the mesh used has been less refined.

TABLE II. COMPARISON BETWEEN THIS PAPER AND LITERATURE OF TENSILE TEST'S PEAK VALUES

Ansys vs FEniCS comparison		
Reference	Peak Value (kN)	Relative Error ^a (%)
This Paper	0.6880	-
[12]	0.7162	3.93
[13] (Isotropic)	0.6807	1.07
[15] ($\ell=0.015$)	0.6842	0.555
[18] ($\ell=0.011$)	0.6476	6,23

^a. Relative to literature values

2. Three-point bending test

The following problem consists of a three-point bending test. More details can be seen in [15]. The geometry and boundary conditions are given in Fig. 8. The mesh was built using the Gmsh preprocessor and has 72768 triangular elements, in which a refinement was performed in the center, where the crack is expected to propagate [18]. The material has a modulus of elasticity $E = 20.8\text{GPa}$, a Poisson coefficient $\nu = 0.3$ and a critical energy release rate $G_c = 0.54\text{MPa mm}$. Thus, the Lamé parameters $\lambda = 12000\text{MPa}$ and $\mu = 8000\text{MPa}$ were obtained. A value of 0.03mm was also used for the characteristic length ℓ , similar to that used in the literature [13,15].

We start with the same numerical modeling for the traction problem but now with a point force. The point (0,0) has zero nodal displacements in x and y . At the point (8,0), the shift is restricted to y only. The force is applied punctually at (4,2). An initial displacement of 0.005mm was defined at the top of the geometry, changing to 0.00001mm when approaching the failure and returning to 0.005mm after the failure to follow the crack in greater detail. Fig. 9 shows the vector field of the phase field variable as a ϕ function of the load level, with representation like the previous problem. The force-displacement curve is shown in Fig. 10. The results obtained are satisfactory and close to the literature, as shown in Table III, even with a variation of displacements different from that used by the authors to reduce computational costs [13,15,18].

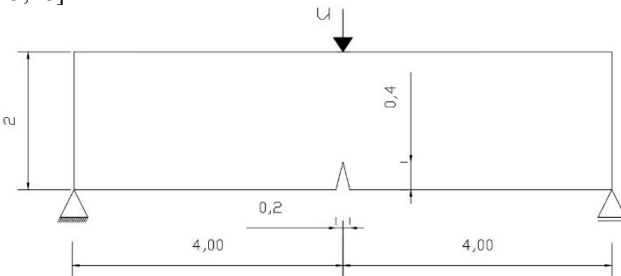


Fig. 8. Geometry and boundary conditions for three-point bending test, adapted from [13]

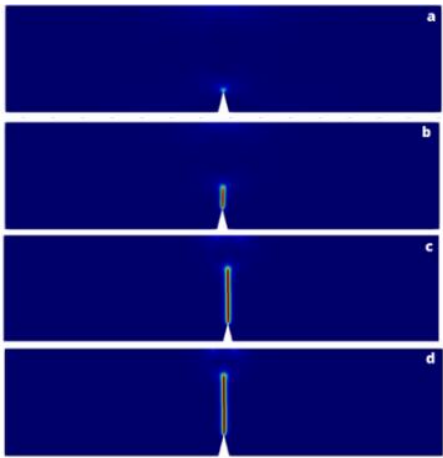


Fig. 9. Phase field for displacements $u=0.04\text{ mm}$ (a), $u=0.045\text{ mm}$ (b), $u=0.056\text{ mm}$ (c), $u=0.071\text{ mm}$ (d)

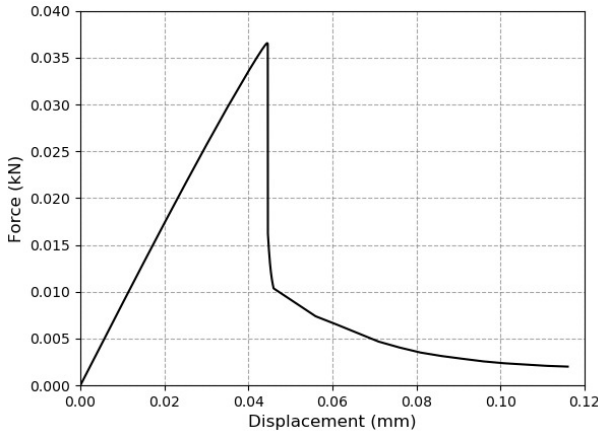


Fig. 10. Force-displacement curve for three-point bending test

TABLE III. COMPARISON BETWEEN THIS PAPER AND LITERATURE OF THREE-POINT BENDING TEST PEAK VALUES

Ansys vs FEniCS comparison		
Reference	Peak Value (kN)	Relative Error ^a (%)
This Paper	0,0365	-
[13] (Hybrid)	0,0417	12,47
[15]	0,0385	5,19
[18]	0,0412	11,4

a. Relative to literature values

IV. CONCLUSIONS

This work proposes the implementation of classic linear elastic fracture mechanics problems based on a tutorial and examples found in the literature, with the help of Gmsh and Paraview. The comparison between the FEniCS and ANSYS software for the elastic linear problem obtained a satisfactory result, with an error, on average, of less than 0.6%. The traction problem presented an error, on average, of approximately 2.9% for the peak value when compared to the literature. The three-point bending problem, on the other hand, presented an error, on average, of approximately 9.7% for the peak value in comparison with the same source. Thus, it is concluded that FEniCS can be used for academic purposes both for solving classic problems of Strength of Materials and MFLE. Furthermore, this tool, together with Gmsh and Paraview, provides users with advanced approaches to learning engineering problem-solving.

ACKNOWLEDGEMENTS

The work was developed under the 2021 Researcher Productivity Notice, and the authors would like to thank Ifes for the opportunity and infrastructure. To CNPq, for the scholarship, allowing full dedication to Scientific Initiation. To FAPES and CNPq - Support Program for Emerging Centers - PRONEM - Notice 06/2019 and Grant Term - 502/2020 - SIAFEM: 220-765DF.

REFERENCES

[1] T. R. Chandrupatla and A. D. Belegundu, "Elementos Finitos," 4^o ed., Pearson, 2014.
[2] Ansys ® [Ansys Student], Version 2022, Ansys, Inc.
[3] M. S. Alnaes et al., "The FEniCS Project Ver. 1.5," Archive of Numerical Software, vol. 3, 2015.

[4] T. L. Anderson, "Fracture Mechanics - Fundamentals and Applications," 2nd Edition, CRC Press, 1994.
[5] K. Majidzadeh, E. M. Kauffmann, and D. V. Ramsamooj, "Application of Fracture Mechanics in the Analysis of Pavement Fatigue," in Association of Asphalt Paving Technologists Proceedings, pp. 227-246, 1971.
[6] G. Clerc, A. J. Brunner, P. Niemi, and J. Kuilen, "Application of Fracture Mechanics to Engineering Design of Complex Structures," in 1st Virtual European Conference on Fracture, Procedia Structural Integrity, vol. 28, pp. 1761-1767, 2020.
[7] G. J. Neate, G. M. Sparkes, H. D. Williams, and A. T. Stewart, "Application of Fracture Mechanics to Industrial Problems," in Fracture Mechanics, Pergamon, pp. 69-90, 1979.
[8] J. A. L. Rocha, "Termodinâmica da Fratura: Uma Nova Abordagem do Problema da Fratura nos Sólidos," EDUFBA, 2010.
[9] P. T. R. Mendonça and E. A. Fancello, "O Método dos Elementos Finitos Aplicado à Mecânica dos Sólidos," 1st Edition, Orsa Maggiore, 2019.
[10] Ansys ®, "Simulation Is a Superpower," Available: <https://www.ansys.com/company-information/simulation-is-a-superpower>. [Accessed: May. 14, 2023].
[11] H. P. Langtangen and A. Logg, "Solving PDEs in Python: The FEniCS Tutorial 1," Simula Research Laboratory and Department of Informatics, University of Oslo, 2017.
[12] E. M. Pañeda and H. Hirshikesh, "Phase Field Fracture Implementation in FEniCS," ResearchGate, 2020.
[13] M. Ambati, T. Gerasimov, and L. De Lorenzis, "A Review on Phase-Field Brittle Models of Fracture and a New Fast Hybrid Formulation," Comput Mech, vol. 55, pp. 383-405, 2015.
[14] J. Amor, J. Marigo, and C. Maurini, "Regularized Formulation of the Variational Brittle Fracture with Unilateral Contact: Numerical Experiments," Journal of the Mechanics and Physics of Solids, vol. 57, pp. 1209-1229, 2009.
[15] C. Miehe, M. Hofacker, and F. Welschinger, "A Phase Field Model for Rate-Independent Crack Propagation: Robust Algorithmic Implementation Based on Operator Splits," Comput Methods Appl Mech Eng, vol. 199, pp. 2765-2778, 2010.
[16] C. Geuzaine and J. Remacle, "Gmsh: A 3-D Finite Element Mesh Generator with Built-in Pre- and Post-Processing Facilities," International Journal for Numerical Methods in Engineering, vol. 79, pp. 1309-1331, 2009.
[17] J. Ahrens, B. Geveci, and C. Law, "Paraview: An End-user Tool for Large Data Visualization," in Visualization Handbook, Elsevier, pp. 978-0123875822, 2005.
[18] H. Hirshikesh, S. Natarajan and R. K. Annabattula, "A FEniCS Implementation of the Phase Field Method for Quasi-Static Brittle Fracture," Frontiers of Structural and Civil Engineering, vol. 13, 2019.

AUTHORS



Caio Prates Martins

I am Caio César Prates Martins. I am 21 years old and live in Espírito Santo - Brazil, with my girlfriend and my mother. I like to study, learn new things, cook, gym, and hang out with my girlfriend and friends.

I am a Metallurgical Engineering student at the Federal Institute of Espírito Santo and I am in the 9th period of the course. I have also been an intern at Biancogres for 3 months. Since I was a kid I like soccer and I support Flamengo. At the age of 14 I left my home state and ventured to São Paulo in search of my dream of becoming a soccer player. I learned a lot about interpersonal relationships and how to manage on my own. Unfortunately, I was not successful, but life went on. At the age of 17 I entered college. Recently I did a scientific initiation that introduced me to Research and Development, helping me to advance my knowledge in Python programming, finite element modeling, and fracture mechanics. At the company where I work, we produce porcelain tiles, ceramic and vinyl tiles with excellence and quality, and I am happy to contribute to this.



Brenda Lemos

I'm Brenda Resende Lemos. I'm 23 years old and I live in Espírito Santo - Brazil, with my parents and my brother. I like to study, play handball, watch football games, go to church and hang out with friends.

At the age of 14, I won a scholarship to play handball at a school that is a reference in the sport. That same year, the team in which I participated won the regional, state and Brazilian. It was an amazing experience, I learned about perseverance and that we should never give up on our dreams. In addition, I am in love with Flamengo and I collect cups and shirts for the team.

I'm a Civil Engineering student at the Federal Institute of Espírito Santo and I'm in the 7th period of the course. I did a scientific initiation on numerical modeling and I'm currently doing one on the use of coconut fiber in cement matrices. I like it and I am very interested in following in the area of research.



Jean de Oliveira Bone

Born in Vitória, Espírito Santo state, has a degree in Mechanical Engineering from Ufes - Federal University of Espírito Santo - ES, completed in 2011. He is currently pursuing his Master's degree in Metallurgical and Materials Engineering at Ifes - Federal Institute of Espírito Santo at the Victory campus, where he researches the influence of welding processes on the mechanical properties of steels and teaches disciplines in the area of Materials Science and Welding at the Guarapari campus of Ifes. Has courses in the areas of non-destructive testing by Abend - Brazilian Association of Non-Destructive Testing and in the area of defect diagnosis by vibration analysis by Fupai - Foundation for Research and Advice to Industry of Unifei - Federal University of Itajubá - MG. He is interested in the areas of welding, heat treatment, fracture mechanics, and materials science. He also has a background in other areas such as post-graduation in Education from Faculdade Pitágoras - MG and a degree in Sociology from Unicesumar - PR, to be concluded in 2023.



André Galdino

Possui graduação em Engenharia de Materiais pela Universidade Federal da Paraíba (1997), mestrado em Engenharia e Ciência de Materiais pela Universidade Federal do Ceará (2003) e doutorado em Engenharia Mecânica pela Universidade Estadual de Campinas (2011). Atualmente é professor no Instituto Federal do Espírito Santo, Campus Vitória, Coordenadoria de Mecânica. Tem experiência na área de Engenharia de Materiais e Metalúrgica, com ênfase em Materiais Metálicos e Cerâmicos, atuando principalmente nos seguintes temas: cerâmicas porosas, propriedades físicas e mecânicas, cerâmica vermelha, metalurgia do pó, compósitos, biomateriais, cerâmicas refratárias e ensino de Materiais. Além disso, é líder nos grupos de pesquisa "Materiais e Processos de Fabricação" e "Materiais de Construção Civil" no Campus Vitória. É professor associado do Instituto Nacional de Ciência e Tecnologia em Biofabricação BIOFABRIS e pesquisador no Polo de Inovação EMBRAPA IFES em Materiais e Metalurgia. É professor permanente no Programa de Pós-Graduação em Engenharia Metalúrgica e de Materiais e professor permanente no Mestrado Profissional em Tecnologias Sustentáveis.

AUTHORS



Rodolfo Mendes

Possui doutorado em Estruturas e Materiais pela COPPE/ Universidade Federal do Rio de Janeiro (2020), mestrado em Estruturas pela Escola Politécnica da Universidade de São Paulo (2012) e graduação em Engenharia Civil pela Universidade Federal do Espírito Santo (2009). Atualmente, é professor de ensino básico, técnico e tecnológico do Instituto Federal de Educação Ciência e Tecnologia do Espírito Santo, Campus Vitória. Tem experiência na área de Engenharia Civil, com ênfase em Estruturas, nos seguintes temas: estruturas de concreto, modelagem numérica do concreto e compósitos cimentícios reforçados com fibras de aço

Essentially non-oscillatory schemes applied to Buckley-Leverett equation with diffusive term

ARTICLE HISTORY

Received 15 February 2023
Accepted 12 May 2023

Raphael de Oliveira Garcia
Federal University of São Paulo
Osasco, Brazil
rogarcia@unifesp.br
ORCID: 0000-0003-2915-150X

Graciele Paraguaia Silveira
Federal University of São Carlos
Sorocaba, Brazil
graciele@ufscar.br
ORCID: 0000-0002-8610-4841

Essentially non-oscillatory schemes applied to Buckley-Leverett equation with diffusive term

Esquemas esencialmente no oscilatorios aplicados a la ecuación de Buckley-Leverett con término difusivo

Raphael de Oliveira Garcia
Actuarial Science Department
Federal University of São Paulo
Osasco, Brazil
rogarcia@unifesp.br
ORCID: 0000-0003-2915-150X

Graciele Paraguaia Silveira
Mathematics, Physics and Chemistry
Department
Federal University of São Carlos
Sorocaba, Brazil
graciele@ufscar.br
ORCID: 0000-0002-8610-4841

Abstract— The purpose of this work was to investigate the flow of two-phase fluids via the Buckley-Leverett equation, corresponding to three types of scenarios applied in oil extraction, including a diffusive term. For this, a weighted essentially non-oscillatory scheme, a Runge-Kutta method and a central finite difference were computationally implemented. In addition, a numerical study related to the precision order and stability was performed. The use of these methods made it possible to obtain numerical solutions without oscillations and without excessive numerical dissipation, sufficient to assist in the understanding of the mixing profiles of saturated water and petroleum fluids, inside pipelines filled with porous material, in addition to allowing the investigation of the impact of adding the diffusive term in the original equation.

Resumen—El propósito de este trabajo fue investigar el flujo de fluidos bifásicos a través de la ecuación de Buckley-Leverett, correspondiente a tres tipos de escenarios aplicados en la extracción de petróleo, incluyendo un término difusivo. Para ello, se implementaron computacionalmente: un esquema esencialmente no oscilatorio ponderado, un método de Runge-Kutta y un esquema de diferencias finitas centrado. Además, se realizó un estudio numérico relacionado con el orden de precisión y estabilidad. El uso de estos métodos permitió obtener soluciones numéricas sin oscilaciones y sin disipación numérica excesiva, suficientes para auxiliar en la comprensión de los perfiles de mezcla de agua saturada y fluidos derivados del petróleo, en el interior de tuberías llenas de material poroso, además de permitir la investigación del impacto de sumar el término difusivo en la ecuación original.

Keywords—Numerical methods, Immiscible two-phase fluid, Petroleum flow

Palabras clave—Métodos numéricos, Fluido bifásico inmiscible, Flujo de petróleo

I. INTRODUCTION

The technological advances achieved after the Industrial Revolution, followed by the development of equipment dependent on energy sources from petroleum, made the oil industry reach a crucial role in the world economy. As it is an exhaustible natural source and a potential pollutant, if disposed of in an erroneous way, the efficiency in the extraction process and the detailed understanding of the phenomena involved become indispensable for the challenges of delivering a product inserted within a clean economy.

A central problem in this sector is the displacement of petroleum through pipes filled with a porous medium, characterized by the injection of another fluid (saturated water) to help maintain the flow inside the tube. In this sense, a mathematical model used to describe the flow of two-phase incompressible fluids is the classical Buckley-Leverett equation [1]. In this work, the authors obtained solutions considered unrealistic due to the appearance of discontinuities, indicating the need for studies on the propagation of singularities.

Since then, different researchers started to investigate numerical solutions to the original Buckley-Leverett flow equation. The characteristic of non-linearity of the partial differential equation enables the use of numerical methods and computational techniques, in an attempt to find approximate solutions without spurious oscillations or excessive numerical dissipation.

Fayers and Sheldon [2] used a finite difference scheme for the partial differential equation for flow two-phase incompressible fluids. In [3], numerical methods dependent on Riemann solvers were applied to follow the discontinuities evolution (shock waves). However, this procedure requires the calculation of propagation velocities at each point of the established grid.

Currently, experiments in laboratory of two-phase flow in porous medium reveal complex profiles that include fluid infiltrations modelled by diffusive and dispersive terms, which suggest modifications to the classical Buckley-Leverett equation [4].

Central schemes of finite volume were used by [5] and a fully space-time mixed hybrid finite element/volume discretization was developed by [6], both to solve the modified Buckley-Leverett equation.

In preliminary studies, Garcia and Silveira [7] investigated a fifth-order weighted essentially non-oscillatory scheme applied in the classical Buckley-Leverett equation. Continuing the research, in this work the authors began to evaluate the addition of a diffusive term.

In this way, the objective of this paper was to apply a weighted essentially non-oscillatory scheme, coupled to a three-stage Runge-Kutta method and a central finite difference scheme in the discretization of the Buckley-Leverett equation with diffusive term and obtain numerical solutions capable of representing the temporal evolution starting from three initial scenarios, represented by discontinuous functions such as Heaviside step, barrier and well rectangular. In addition to carrying out a study on the impact of the diffusive term on the numerical solution of the classical Buckley-Leverett equation. The modeling of the phenomenon under study will be explained in Section 2, the numerical methods used are in Section 3 and, finally, in Section 4, the results achieved regarding the evolution of the mixture between saturated water and oil are shown. The computational implementation was carried out in *Octave*, with our own codes.

II. MATHEMATICAL MODELING

Currently, a problem of global interest is the extraction of petroleum underground through a tube filled with a porous medium. After drilling the soil to the underground oil reservoir, a certain amount is drained due to the high pressure that the oil is found, but as the extraction progresses, there is a decrease in pressure with consequent interruption of the flow, still leaving a lot of petroleum in the subsoil. A standard method subsequent to the initial extraction is to pump water into rest oil reservoir to force the continuation of extraction. In this case, the fluid is two-phase, oil and water, and the flow is restarted in the porous medium consisting of rock or sand.

The mathematical modeling consists of representing the oil one-dimensional flow through the tube filled with porous material, by water pumping [8]. Such a model was initially proposed by Buckley and Leverett, 1942 [1], in studies on the flow of two-phase incompressible fluids in porous media.

Let $0 \leq q(x, t) \leq 1$ be the fraction of saturated water and $1 - q(x, t)$ the fraction of oil contained in a pipe filled with a porous material. Such fluids are essentially incompressible, which ensures that the total flow between the pipe ends is equal to any smaller portion of the pipe.

In this way, in regions of the tube where $q = 0$ (pure oil) or $q = 1$ (pure water), the velocities are constant and distinct, but when $0 < q < 1$, the difference between the surface tensions of fluids causes them to move and mix. Buckley and Leverett proposed a model in which the rate of change of q over time (q_t) is described by the following conservation law:

$$q_t + f(q)_x = 0, \quad (1)$$

in which $f(q) = \frac{q^2}{q^2 + a^2(1+q^2)}$ is the water flux, $0 < a < 1$ represents the porosity of the medium and $1 - f(q)$ is the oil flux, with $q = q(x, t)$. Equation (1) models a flow from left to right, in which the tube thickness does not influence the dynamics in question.

The reestablishment of the oil flow, from left to right, can be done by filling part of the pipe on the left with saturated water, allowing the resumption of oil extraction. An initial condition for modeling that procedure is the following conservation law:

$$q(x, 0) = \begin{cases} 1, & \text{if } x < 0 \\ 0, & \text{if } x > 0 \end{cases} \text{ (scenario 1)}. \quad (2)$$

Another situation that can occur is, after an injection of water, the flow is restarted and a second oil portion enters the pipeline, soon after the interruption in the supply of saturated water. This is the desired circumstance, when only an amount of water injected is enough for the continuity of the extraction. The condition to represent that scenario is the conservation law:

$$q(x, 0) = \begin{cases} 0, & \text{if } x < a \\ 1, & \text{if } a < x < b \text{ (scenario 2).} \\ 0, & \text{if } x > b \end{cases} \quad (3)$$

The third situation considered arises when following the injection of water, a small amount of oil is extracted and then there is a new interruption of flow, making it necessary to inject more saturated water after a small amount of oil is extracted. So,

$$q(x, 0) = \begin{cases} 1, & \text{if } x < a \\ 0, & \text{if } a < x < b \text{ (scenario 3).} \\ 1, & \text{if } x > b \end{cases} \quad (4)$$

The three situations represented mathematically by the initial conditions, in (2)-(4), have the flow described by (1).

Considering the spatial domain $[x_i, x_f]$, the boundary condition on the left was of the *Dirichlet type* and on the right of the *radiation type*, that is, when the dynamics reach the boundary, the flow simply goes through the boundary, is not being affected by the edge.

Adding the effects of infiltration in porous media, van Duijn, Peletier, and Pop [9] proposed the following modification to (1):

$$q_t + f(q)_x = \varepsilon q_{xx} + \varepsilon^2 \kappa q_{xxt}, \quad (5)$$

in which ε is a diffusibility coefficient and κ is a dispersive coefficient. As one of the objectives of this work is to observe the effect of diffusion in the Buckley-Leverett equation (1), it was considered that $\varepsilon \gg \varepsilon^2 \kappa$, so the dispersive term was neglected.

The choice of numerical methods to solve (1) and (5) is essential, as a numerical scheme that has a high degree of

numerical dissipation can stand out and mask the real effect of the diffusive term of (5), impairing the interpretation of the results.

III. NUMERICAL METHODS

Choosing a numerical method requires care that depends on the differential equation and the initial conditions imposed by the model under study. The discontinuities of the functions defined by (2), (3) and (4) can evolve to excessive dispersion (oscillation) and/or numerical dissipation, harming the quality of the numerical solution. Thus, for this work, a weighted essentially non-oscillatory scheme (WENO-5 method) was adopted for spatial discretization and the third-order Runge-Kutta TVD method (Total Variation Diminishing) was assumed for temporal discretization.

A. Non-oscillatory schemes and WENO-5

The essentially non-oscillatory (ENO) schemes proposed by [10] and [11] proved to be efficient in terms of decreasing numerical dissipation, in addition to avoiding oscillations in the discontinuity regions of the solutions. One important properties of these methods is to determine the smoothest stencil among the options, in order to preserve high order of accurate.

In general, to approximate the flow a r th-order ENO scheme selects the smoothest stencil among r possibilities [12]. On the other hand, the weighted essentially non-oscillatory (WENO) schemes make a convex combination of all stencils, for the numerical flow approximation. A weight is designated to each stencil, describing its contribution portion to the process.

The weights are defined by optimal weights in smooth regions, while maintaining the high order of accuracy. In non-smooth regions, weights close to zero are assigned for stencils that contain discontinuities [13].

The WENO method technique is based on the flow definition of the ENO schemes, considering a one-dimensional conservation law $u_t + f(u)_x = 0$. The spatial operator that approximates $-f(u)_x$ in x_j is

$$L = -\frac{1}{\Delta x}(\hat{f}_{j+1/2} - \hat{f}_{j-1/2}), \quad (6)$$

in which Δx is the spatial discretization size and $\hat{f}_l$ is the numerical flux [13].

The r stencils are denoted by S_k , $k = 0, 1, \dots, r-1$, of the form conservation law:

$$S_k = \{x_{j+k-r+1}, x_{j+k-r+2}, \dots, x_{j+k}\},$$

in which defines the amount of points of S_k , that will be applied to calculate the value of $\hat{f}_{j+1/2}$. Therefore, a 3rd-order ENO scheme ($r = 3$) is going to have: S_k , $k = 0, 1, 2$, described by $S_k = \{x_{j+k-2}, x_{j+k-1}, x_{j+k}\}$, which results in $S_0 = \{x_{j-2}, x_{j-1}, x_{j+k}\}$ for $k=0$, $S_1 = \{x_{j-1}, x_j, x_{j+1}\}$ for $k=1$, and $S_2 = \{x_j, x_{j+1}, x_{j+2}\}$ for $k=2$.

ENO schemes approximate $\hat{f}_{j+1/2}$ through a polynomial interpolation at the points of each stencil [13] and [14]. This approximation is given by

$$\hat{f}_{j+1/2} = q_k^r(f_{j+k-r+1}, f_{j+k-r+2}, \dots, f_{j+k}), \quad (7)$$

in which

$$q_k^r(g_0, \dots, g_{r-1}) = \sum_{l=0}^{r-1} a_{k,l}^r g_l.$$

Let $g(x)$ be a smooth function. The average approximation of $g(x)$ in cell I_j is defined by:

$$\bar{g}_j = \frac{1}{\Delta x} \int_{x_{j-1/2}}^{x_{j+1/2}} g(\xi) d\xi,$$

where $\xi \in I_j = (x_{j-1/2}, x_{j+1/2})$.

To obtain the constants $a_{k,l}^r$ in (7), consider the primitive function of $g(x)$ defined by $G(x) = \int_{-\infty}^x g(\xi) d\xi$. The value of $G(x_{j+1/2})$ is:

$$G(x_{j+1/2}) = \sum_{i=-\infty}^j \int_{x_{j-1/2}}^{x_{j+1/2}} g(\xi) d\xi = \sum_{i=-\infty}^j \bar{g}_i \Delta x_i.$$

It means that, once the average approximations of the cells $\bar{g}_i$ are known, then $G(x)$ at the boundary of cell I_j are also known. Thus, the constants $a_{k,l}^r$ are determined by interpolating $G(x_{j+1/2})$ by a r degree polynomial $P(x)$, at most. Therefore, approximation is given by:

$$\hat{f}_{j+1/2} = p(x_{j+1/2}) = P'(x_{j+1/2}) = \sum_{l=0}^{r-1} a_{k,l}^r \bar{g}_l, \quad (8)$$

in which $a_{k,l}^r$ are obtained from the Lagrange interpolating polynomial [14], with the data in Table I.

TABLE I. $a_{k,l}^r$ COEFFICIENTS

r	k	$l = 0$	$l = 1$	$l = 2$
3	0	1/3	-7/6	11/6
	1	-1/6	5/6	1/3
	2	1/3	5/6	-1/6

A r th-order ENO scheme leads to a $(2r-1)$ th-order WENO scheme, so a 3rd-order ENO results to a 5th-order WENO. As mentioned, in the WENO method, for each possible stencil S_k , $k = 0, 1, \dots, r-1$, a weight ω_k is assigned and these are used to calculate the numerical flux:

$$\hat{f}_{j+1/2} = \sum_{k=0}^{r-1} \omega_k q_k^r(f_{j+k-r+1}, f_{j+k-r+2}, \dots, f_{j+k}). \quad (9)$$

The weight ω_k for the stencil S_k is defined by given

$$\omega_k = \frac{\alpha_k}{\alpha_0 + \dots + \alpha_{r-1}}, \text{ with } \alpha_k = \frac{C_k^r}{(\varepsilon + IS_k)^p},$$

and $k = 0, 1, \dots, r-1$.

Taking $p = r$, the coefficients C_k^r are optimal values to determine ω_k [14]. The term IS_k is an indicator of smoothness and for $r = 3$ we have:

$$\begin{aligned} IS_0 &= \frac{13}{12}(f_{j-2} - 2f_{j-1} + f_j)^2 + \frac{1}{4}(f_{j-2} - 4f_{j-1} + 3f_j)^2; \\ IS_1 &= \frac{13}{12}(f_{j-1} - 2f_j + f_{j+1})^2 + \frac{1}{4}(f_{j-1} - f_{j+1})^2; \\ IS_2 &= \frac{13}{12}(f_j - 2f_{j+1} + f_{j+2})^2 + \frac{1}{4}(3f_j - 4f_{j+1} + f_{j+2})^2. \end{aligned}$$

This measure was introduced by [13], with the aim of achieving high accurate for the case where $r = 3$. Note that as IS_k increases, the smoothness decreases and, consequently, α_k becomes close to zero as does ω_k , meaning that a weight close to zero will be assigned to non-smooth solutions.

B. Third-order Runge-Kutta TVD

Once the spatial discretization is concluded, a method for temporal discretization that maintains the non-oscillatory characteristics achieved is necessary.

Numerical methods belonging to the TVD class (Total Variation Diminishing) have the property of avoiding oscillations that are not typical of the phenomenon under study [8]. A good alternative is the high-order Runge-Kutta TVD methods, which were developed by [12] in research related to efficient implementations for ENO's schemes.

A Total Variation Diminishing (TVD) technique has the following definition: if, for any data set U^n , the values U^{n+1} computed satisfy $TV(U^{n+1}) \leq TV(U^n)$, where

$$TV(U^n) = \sum_{i=1}^N |U_i^n - U_{i-1}^n|,$$

is the total variation. In this work, a third-order Runge-Kutta TVD (RK3-TVD) method was chosen, whose expressions are given by

$$\begin{aligned} u^{(1)} &= u^n + \Delta t L(u^n) \\ u^{(2)} &= \frac{3}{4}u^n + \frac{1}{4}u^{(1)} + \frac{1}{4}\Delta t L(u^{(1)}), \\ u^{n+1} &= \frac{1}{3}u^n + \frac{2}{3}u^{(2)} + \frac{2}{3}\Delta t L(u^{(2)}) \end{aligned} \quad (10)$$

in which L is the spatial operator of differential equation.

The WENO and RK3-TVD methods numerically solve the classical Buckley-Leverett equation. To discretize the diffusive term of the modified Buckley-Leverett equation, we use a central finite difference scheme.

C. Fourth-order central finite difference scheme

Once there are at least two ways to numerically add the q_{xx} term in the conservation law: one is to discretize the diffusive term in the flux context [15], given by:

$$q_{xx}(x_{i+1/2}, t) \cong \frac{\bar{q}_{i+3/2} - 2\bar{q}_{i+1/2} + \bar{q}_{i-1/2}}{\Delta x^2}, \quad (11)$$

and the other in the finite difference context [5], which

$$q_{xx}(x_i, t) \cong \frac{-\bar{q}_{i-2} + 16\bar{q}_{i-1} - 30\bar{q}_i + 16\bar{q}_{i+1} - \bar{q}_{i+2}}{12\Delta x^2}. \quad (12)$$

In this work, both discretizations were performed, as for the interpretation of the results there were no significant differences between them, we selected (12) to keep in our codes, is that, the fourth-order central finite difference scheme (CFDS-4). We use a central finite difference scheme.

D. Numerical study of the methods

The numerical study of the WENO-5 method makes it possible to verify its convergence order (which is the number of significant algarisms that are correct when the solution is obtained), that is, it allows the computational verification of the consistency of the method. For numerical simulations, consider the function, $f(x) = \sin(x)$, for $x \in [0, \pi]$ and the values for the spacing $h = \pi 2^{-p}$, $6 \leq p \leq 11$.

Asymptotically, the error made by the approximation has a behavior of the form $E = \max|\text{analytical solution} - \text{numerical solution}| = Ch^\alpha$, where C is a positive constant independent of h and α is the order of convergence. Thus, in the *dilog* graph, $\log(E(h)) \times \log(h)$, the error has an approximately linear behavior and the slope of the line provides the order α of the method. Fig. 1 shows the scatter plot for the fit of the model, for calculating the order of convergence of the WENO-5 method, whose value obtained was $\alpha = 4.96$, that is, a 5th-order.

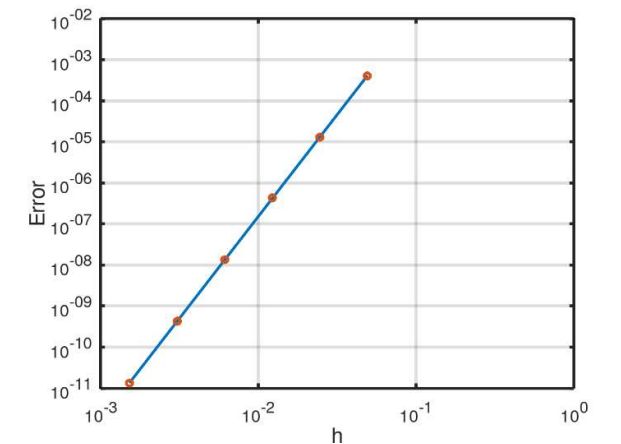


Fig. 1. Scatter plot for calculating the order of convergence of the WENO-5

The numerical study of the RK3-TVD method was carried out with the aim of verifying the order of convergence, that is, computing the consistency of the method. For numerical simulations, we consider the initial value problem $x' = 5t - 2x$, with $x(0) = 1$, $t \in [0, 2]$ and with the spacings $h = 2^{1-p}$, $3 \leq p \leq 10$. The analytical solution is $x(t) = (9e^{-2t} + 10t + 5)/4$.

The error for the approximation is obtained following the same procedures for the numerical study of the WENO-5. Thus, Fig. 2 shows the scatter plot for the fit of the model, for calculating the order of convergence of the RK3-TVD method, whose order obtained was $\alpha = 3.07$ (3rd-order).

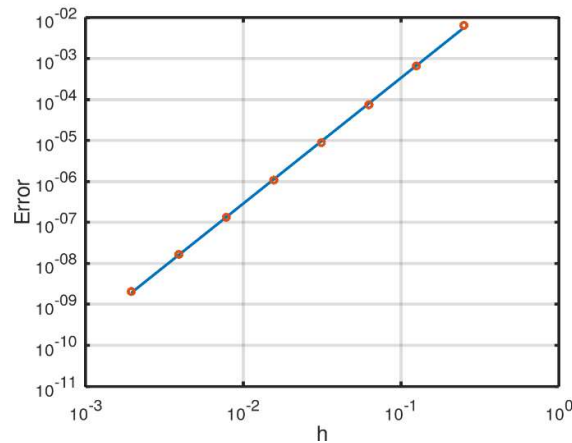


Fig. 2. Scatter plot for calculating the order of convergence of the RK3-TVD

The numerical study of the CFDS-4 allows to obtain consistency of the method, Fig. 3. For numerical simulations, consider the function $f(x) = \sin(x)$, for $x \in [0, 2\pi]$ and the values for the spacing $h = \pi 2^{1-p}$, $3 \leq p \leq 8$.

Estimating the second derivative $f''(x) = -\sin(x)$ of f via CFDS-4, determining the maximum absolute error for each h , and finding the slope on the *dilog* graph, we get $\alpha = 3.99$ (4th-order).

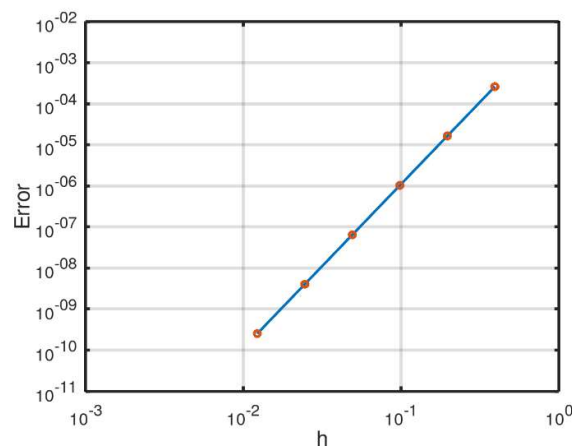


Fig. 3. Scatter plot for calculating the order of convergence of the CFDS-4

E. Analysis of stability

To find a relationship between the spacings Δx and Δt , which keep stable a numerical method obtained by spatial and temporal discretizations via WENO-5 and RK3-TVD, respectively, the linear stability analysis is applied.

The importance of this relationship is to guarantee the convergence of the approximate solution using Lax Equivalence Theorem [16], which states: for any consistent method to generate a sequence of convergent approximate solutions, it is necessary and sufficient that the method be stable, that is, a consistent method is convergent if and only if it is stable.

In this case, the semi-discretization of the solution is represented by the discrete Fourier series in space [16], so by the superposition principle, it is possible to work with only one

term of the series and in this way, the numerical solution can be represented by:

$$u_j(x, t) = \hat{u}_m(t) e^{ij\theta_m}, \theta_m = \omega_m \Delta x, \quad (13)$$

with $m = -N/2, \dots, N/2$.

We consider that the operator L , defined by a conservation law, is written in the form $L(u_{j-r}, \dots, u_{j+s}) = z(\theta_m)$. Temporal semi-discretization is performed using the RK3-TVD method. As it is an explicit scheme, the solution in $t_{n+1} = (n+1)\Delta t$ is represented by $u_j^{n+1} = g(\hat{z}_m) u_j^n$, $\hat{z}_m = -\sigma z(\theta_m)$, $m = -N/2, \dots, N/2$, $n \geq 1$, where g is the amplification factor that depends on θ_m and $\sigma = \Delta t / \Delta x$.

Thus, a spatial discretization coupled with a time discretization will be stable, if the amplification factor satisfies the following discrete von Neumann criterion, $|g(-\sigma z(\theta))| \leq 1$, $\forall \theta \in [0, 2\pi]$. Such a stability condition establishes an upper bound on σ , which keeps the method linearly stable. Therefore, if the domain of the spatial variable was discretized by a regular spacing Δx , the stability condition allows finding a value for the temporal spacing Δt , which keeps the method stable.

WENO-5 linearization and stability analysis

Let L be a spatial operator of a conservation law described by $L(u_{j-3}, \dots, u_{j+2}) = \hat{f}_{j+1/2} - \hat{f}_{j-1/2}$. We linearize $\hat{f}_{j+1/2}$ [17], considering a solution represented by the Fourier series

and we substitute it in the conservation law, we have the amplification factor given by:

$$z(\theta_m) = \frac{16}{15} \sin^6 \frac{\theta_m}{2} + i \left(-\frac{1}{6} \sin 2\theta_m + \frac{4}{3} \sin \theta_m + \frac{16}{15} \sin^5 \frac{\theta_m}{2} \cos \frac{\theta_m}{2} \right), \quad (14)$$

in which $\theta_m \in [0, 2\pi]$. Equation (14) is represented in the complex plane Fig. 4, varying θ_m from 0 to 2π , black line.

RK3-TVD stability

For stability analysis of the RK3-TVD method, the polynomial that characterizes the method is considered, $p(\lambda) = \frac{1}{6}\lambda^3 + \frac{1}{2}\lambda^2 + \lambda + 1$, as the amplification factor $g(z) = \frac{1}{6}z^3 + \frac{1}{2}z^2 + z + 1$, the boundary of the stability region is $\partial S_t = \{z: |g(z)| = 1\}$, where $g(z) = e^{i\phi}$, with $\phi \in [0, 2\pi]$. Hence, for each value of ϕ , a third degree polynomial equation is defined. Solving the equations for each ϕ in the *Octave* and selecting the roots with the highest magnitude for each ϕ , we obtain the region of stability represented by the red color in Fig. 4.

So, the computational implementation of the WENO-5 numerical method with RK3-TVD consists of applying the temporal discretization of the classical Buckley-Leverett equation by (10), with the spatial operator defined by (6) and (7).

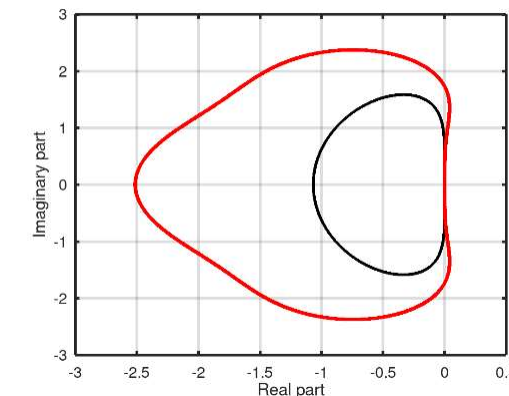


Fig. 4. Stability region of the RK3-TVD method, in red, compared to stability region of the WENO-5 scheme

CFDS-4 stability

With regard to the stability analysis for this method, the value found in the literature for second-order central finite difference schemes applied in the heat equation was used as a reference, that is, $\epsilon \frac{\Delta t}{\Delta x^2} \leq \frac{1}{2}$ [8].

IV. SIMULATIONS

In this section, the simulations for the classical Buckley-Leverett equation (Section IV.A), the modified Buckley-Leverett equation (Section IV.B) and an analysis of diffusive term (Section IV.C) are shown.

A. Classical Buckley-Leverett equation

The simulations were performed for three scenarios, according to the initial conditions defined in Section 2. For all scenarios, we have $x \in [-1, 1]$ with 128 subintervals, $x = 1/64$, and $\Delta t = 0.1\Delta x$, satisfying the stability condition for all methods, according to Fig. 4. Furthermore, $a = 0.5$ was assigned to the constant that characterizes the porous medium, in (1).

The first scenario is represented by the initial condition whose shape is a step, Fig. 5. The numerical solution obtained after 256 iterations is in Fig. 6. The comparison of the graphs reveals the emergence of a region where the fluids mix, between the values of $q = 1$ (pure water) and $q = 0$ (pure petroleum), while the fluid dynamics develops to on the right.

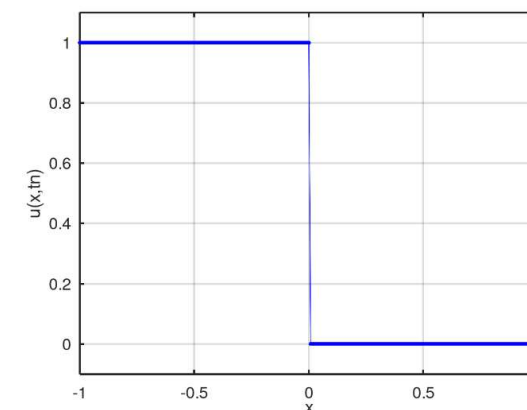
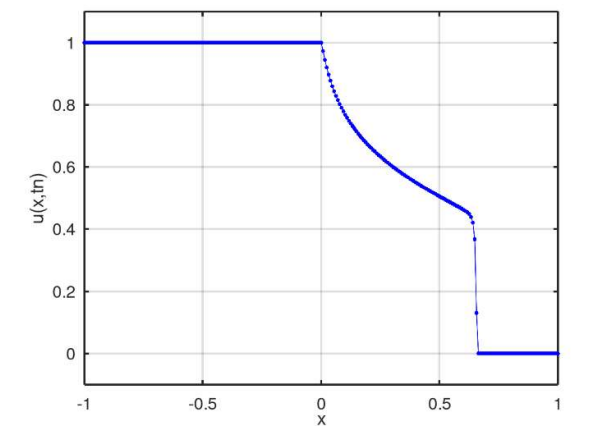


Fig. 5. Initial condition of the scenario 1


 Fig. 6. Numerical solution after 256 iterations, $t_n = 0.4$

The second scenario has the barrier function as an initial condition, Fig. 7. In the numerical solution obtained after 128 iterations, Fig. 8, it is possible to verify two different profiles of mixture between saturated water and petroleum, a linear profile and a non-linear profile similar to the first scenario.

Linear mixing is faster than second mixing and Fig. 9 displays the moment when the linear mixture reaches the non-linear profile. In Fig. 10, the interference of the linear mixture can be noted, that is, there is a decrease in the peak of the graph, with no region containing only water.

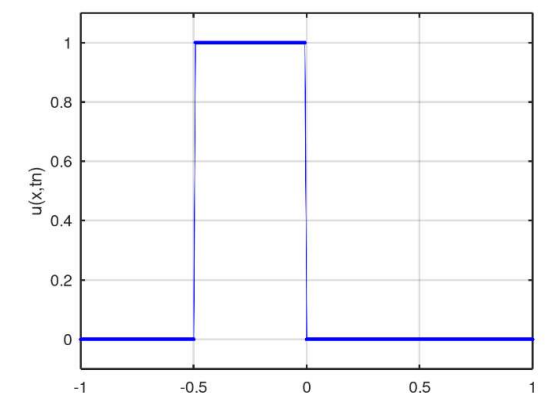
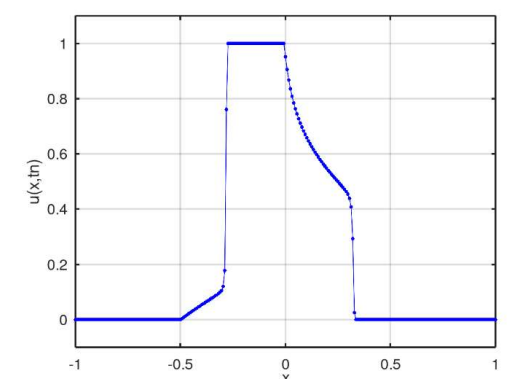
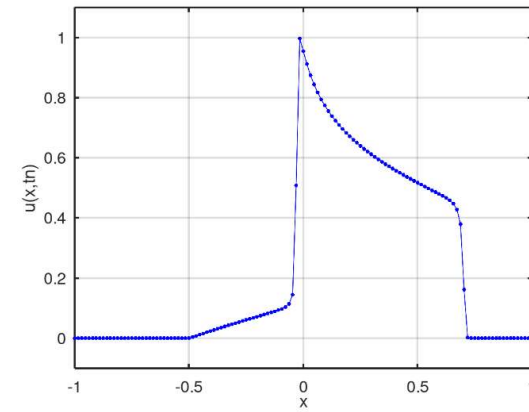
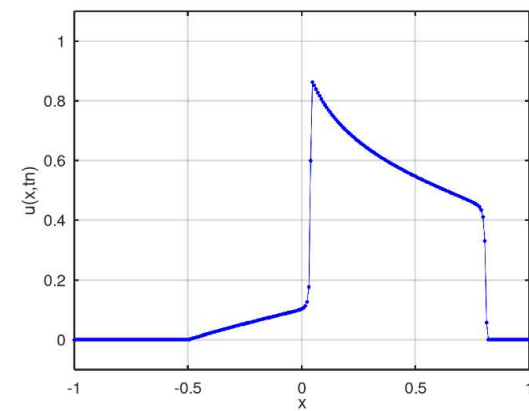


Fig. 7. Initial condition of the scenario 2


 Fig. 8. Numerical solution after 128 iterations, $t_n = 0.2$

Fig. 9. Numerical solution after 278 iterations, $t_n \approx 0.43$ Fig. 10. Numerical solution after 320 iterations, $t_n \approx 0.5$

The third scenario presents an initial condition in the well format, Fig. 11. In the numerical solution found after 128 iterations, two mixing profiles similar to the second scenario is obtained, Fig. 12. However, linear mixing is slower than non-linear mixing, which in turn is reached by the other profile after 175 iterations, Fig. 13. After 398 iterations, there is the moment when the nonlinear profile is close to the fluid region with only water, Fig. 14.

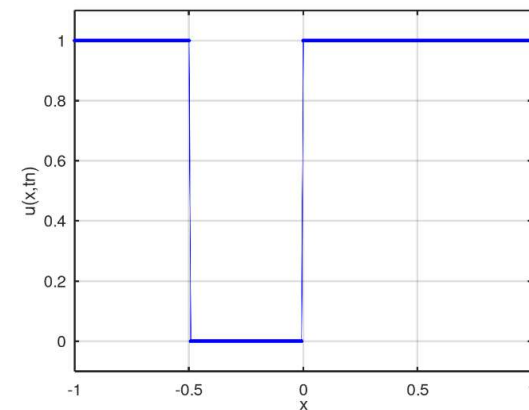
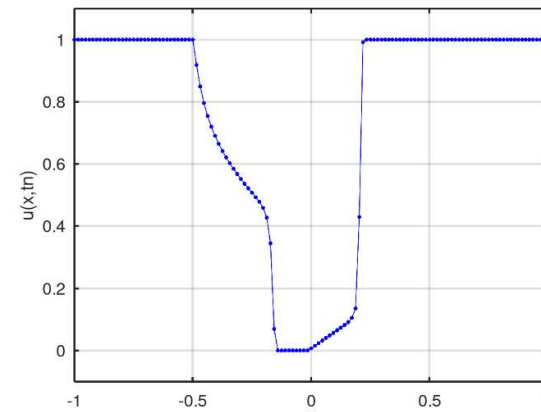
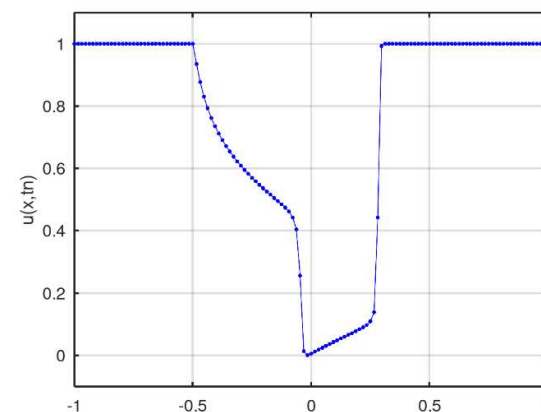
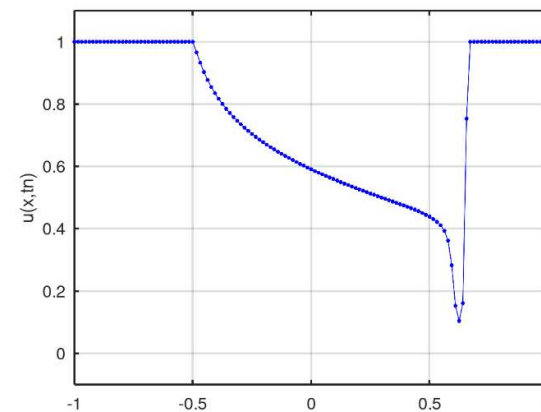


Fig. 11. Initial condition of the scenario 3

Fig. 12. Numerical solution after 128 iterations, $t_n = 0.2$ Fig. 13. Numerical solution after 175 iterations, $t_n \approx 0.27$ Fig. 14. Numerical solution after 398 iterations, $t_n \approx 0.62$

After observing the solution profiles of the classical Buckley-Leverett equation for each scenario, the next step is to consider the addition of the diffusive term and perform comparisons between the solutions.

B. Modified Buckley-Leverett equation

In order to investigate the diffusive term, we considered (5) with $\kappa = 0$, that is, to (1):

$$q_t + f(q)_x = \varepsilon q_{xx}, \quad (15)$$

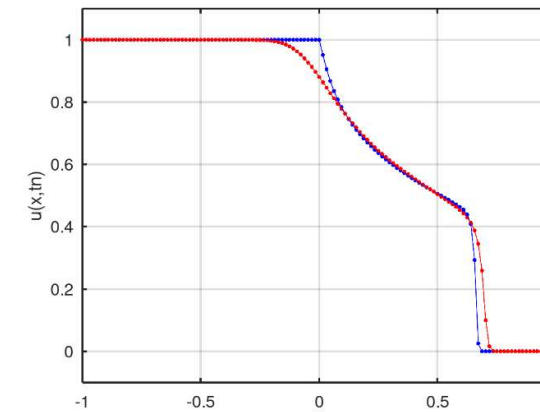
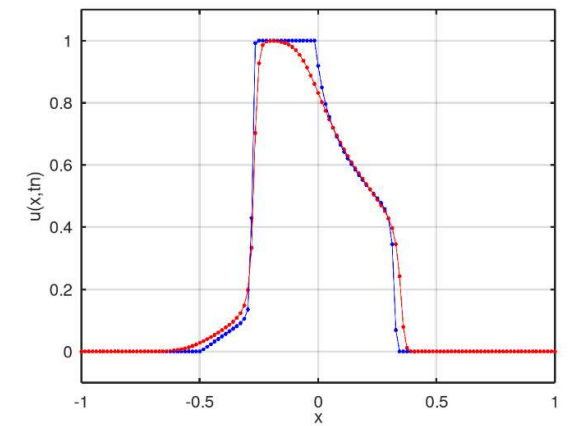
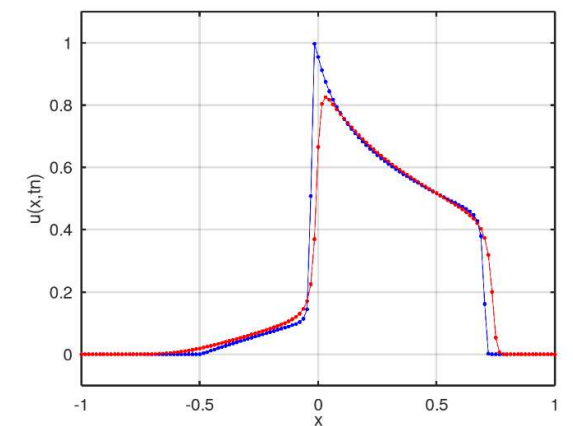
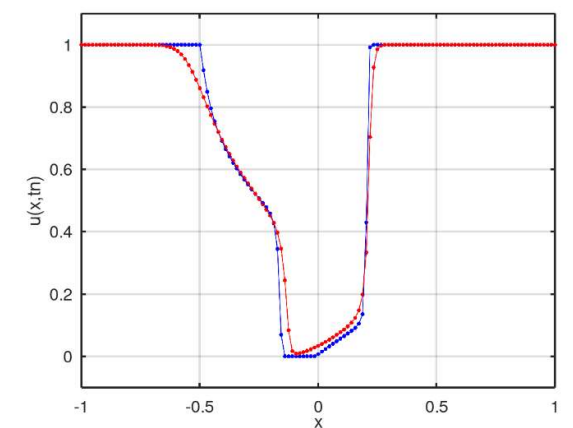
in which $\varepsilon = 0.01$. In the rest, the same scenarios of the Section (IV.A) were adopted, with the same profile of space and time discretizations. For all simulations in this section, blue color represents numerical solution obtained by classical Buckley-Leverett equation and red color represents solutions of the modified Buckley-Leverett equation.

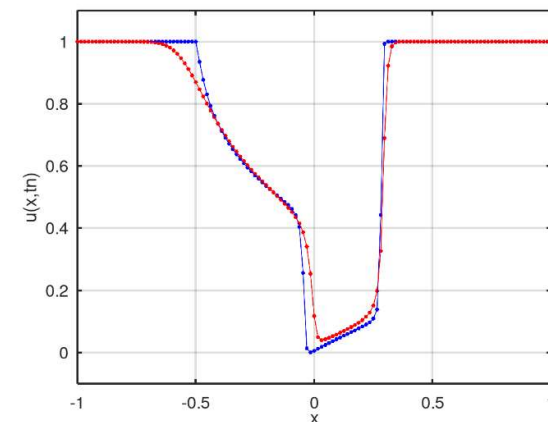
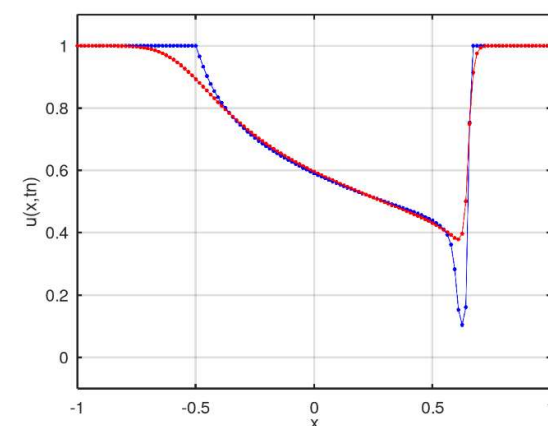
The first scenario with diffusive term (Fig. 15) reveals that there is more water scattered in the oil than in the scenario without diffusion. This is noticeable by the fact that the transitions between pure water, mixture and pure petroleum profiles became smoother.

In terms of execution time, when adding CFDS-4 in the code, the CPU time went from 73.93s to 77.34s, that is, an increase of 4.6%.

The second and third scenarios have similar behavior to the first scenario. The transitions are smoothed by the diffusion, increasing the mixing region between water and oil. See Figs. 16 and 17 for second scenario, and Figs. 18, 19 and 20 for third scenario, where the blue graph represents the numerical solution of the classical Buckley-Leverett equation and the red graph is the solution obtained from the same equation, adding diffusive term.

In Fig. 16, there is an intense smoothing close to $x = 0$ and such smoothing is maintained as time evolves, Fig. 17. This effect is characteristic of the added diffusion in the classical Buckley-Leverett equation.

Fig. 15. Numerical solution after 256 iterations, $t_n = 0.4$ Fig. 16. Numerical solution after 128 iterations, $t_n = 0.2$ Fig. 17. Numerical solution after 278 iterations, $t_n \approx 0.43$ Fig. 18. Numerical solution after 128 iterations, $t_n = 0.2$

Fig. 19. Numerical solution after 175 iterations, $t_n \approx 0.27$ Fig. 20. Numerical solution after 398 iterations, $t_n \approx 0.62$

In the graphs of Figs. 18 and 19, one can see the evolution of two more intense smoothings, one close to $x = -0.5$ and the other close to $x = 0$. In Fig. 20, the smoothings are close to $x = -0.5$ and $x = 0.6$.

Thus, we were able to follow the temporal evolution of the mixture between water and oil in two situations, one where there is only influence of the non-linear advective term (1), and another where there is also a diffusive term (5).

C. Diffusion effect

The focus of the simulations in this section is to verify the impact of the diffusivity coefficient on the solutions of the Buckley-Leverett equation. For this, in the three scenarios presented in Sections (IV.A) and (IV.B), five different values for ε were used, keeping the final time fixed in each scenario.

For all scenarios, the following colors and values were defined: $\varepsilon = 0$ (blue color); $\varepsilon = 0.001$ (green color); $\varepsilon = 0.01$ (red color); $\varepsilon = 0.05$ (black color) and $\varepsilon = 0.1$ (magenta color).

For the first three values of ε , the relationship between the spacings was $\Delta t = 0.1\Delta x$ and for the last two values, it was necessary to define $\Delta t = 0.05\Delta x$ to keep the numerical methods stable.

With the final time of $t_f = 0.2$ for scenario 1, Fig. 21 shows the numerical solution for the five values of ε . It is noticed that the closer the ε value is to zero, the more the solution profile resembles the classic Buckley-Leverett solution, and as the ε value increases, the more the solution is influenced by diffusion, spreading each time plus the mixture of water and oil in the pipeline, represented by the interval $[-1, 1]$.

On the scenario 2, Fig. 22, and scenario 3, Fig. 23, follow the same behavior of the increase in the mixture spreading, as we increase the ε value, observed in scenario 1 (Fig. 21). In both scenarios, the final time was $t_f = 0.2$.

With essentially non-oscillatory methods, which have low numerical dissipation, it was possible to vary the diffusibility coefficient and analyze its impact on the solution of the Buckley-Leverett equation, Figs. 21, 22 and 23.

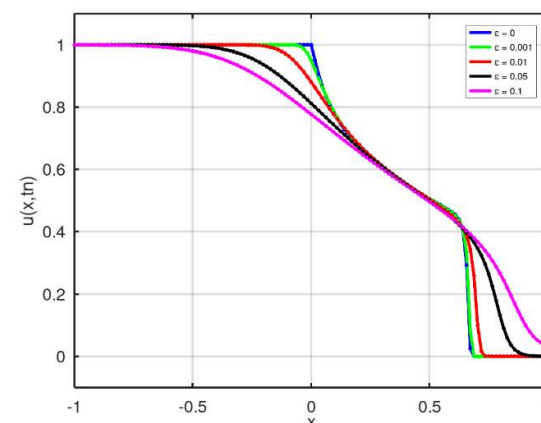


Fig. 21. Diffusion effect on the scenario 1

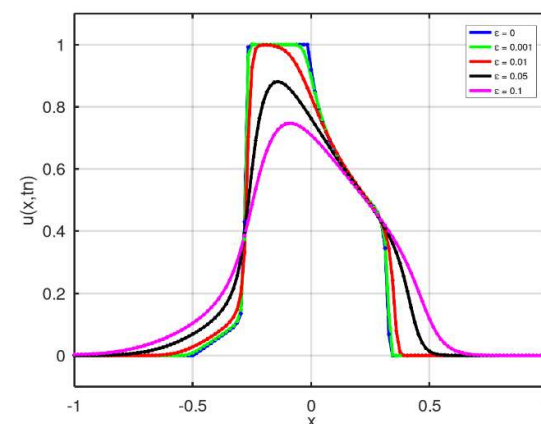


Fig. 22. Diffusion effect on the scenario 2

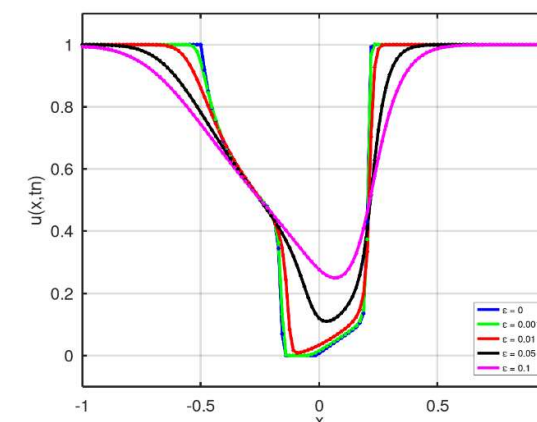


Fig. 23. Diffusion effect on the scenario 3

V. CONCLUSION

In this study, the WENO-5 and RK3-TVD methods were approached, together with a numerical study on the order of convergence of each method, as well as a stability analysis. The methods were applied to the classical Buckley-Leverett equation in order to investigate the temporal evolution of three real scenarios, which may occur during petroleum extraction by injection of saturated water.

The simulations showed mixtures with linear and non-linear profiles due to discontinuities in the initial conditions of each scenario. It is important to note that during the transitions between the profiles there were no oscillations of the numerical solutions or excessive dissipation, indicating that the combination of the WENO-5 and RK3-TVD methods, within the stability condition, provides solutions sufficiently close to the analytical solution.

By adding the diffusive term in the Buckley-Leverett equation, we use the finite difference scheme to discretize it and evaluate the impact of diffusion on the solutions of the three flow scenarios studied. With the methods applied, we were able to vary the parameter ε and observe that the lower its value, the numerical solution approaches the classical Buckley-Leverett solution, and the higher the value of ε , the more the mixture between water and oil is spread through the pipeline, smoothing the transitions between regions containing water, mixture and petroleum.

REFERENCES

- [1] S. E. Buckley, and M. C. Leverett, “Mechanism of fluid displacement in sands,” *Trans. AIME*, vol. 146, pp. 187–196, 1942.
- [2] F. L. Fayers, and J. W. Sheldon, “The effect of capillary pressure and gravity on two-phase fluid flow in porous medium,” *Trans. AIME*, vol. 216, pp. 147–155, 1959.
- [3] W. Proskurowski, “A note on solving the Buckley-Leverett equation in the presence of gravity,” *J. Comput. Phys.*, vol. 41, pp. 136–141, 1981.
- [4] D. A. DiCarlo, “Experimental measurements of saturation overshoot on infiltration,” *Water Resour. Res.*, vol. 40, pp. 4215–1–4215.9, 2004.
- [5] Y. Wang, and C-Y. Kao, “Central schemes for the modified Buckley-Leverett equation,” *J. Comput. Sci.*, vol. 4, pp. 12–23, 2013.
- [6] E. Abreu, and J. Vieira, “Computational numerical solutions of the pseudo-parabolic Buckley-Leverett equation with dynamic capillary pressure,” *Math. Comput. Simulation*, vol. 137, issue C, pp. 29–48, 2017.
- [7] R. O. Garcia, and G. P. Silveira, “Essentially non-oscillatory schemes applied to Buckley-Leverett equation,” in *XXV ENMC - Encontro Nacional de Modelagem Computacional, XIII ECTM - Encontro de Ciência e Tecnologia de Materiais, 9º MCSul - Conferência Sul em Modelagem Computacional, and IX SEMENGO - Seminário e Workshop em Engenharia Oceânica*, 2022 [Online]. Available: <https://www.even3.com.br/anais/enmcsculsemengo2022/534872-essentially-non-oscillatory-schemes-applied-to-buckley-leverett-equation/> [Accessed: Feb. 2, 2023].
- [8] R. J. Leveque, *Finite Volume Methods for Hyperbolic Problems*. New York: Cambridge University Press, 2002.
- [9] C. J. van Duijn, L. A. Peletier, and I. S. Pop, “A new class of entropy solutions of the Buckley-Leverett equation,” *SIAM J. Math. Anal.*, vol. 39, no. 2, pp. 507–536, 2007.
- [10] A. Harten, and S. Osher, “Uniformly high-order accurate non-oscillatory schemes I,” *SIAM J. Numer. Anal.*, vol. 24, pp. 279–309, 1987.
- [11] A. Harten, S. Osher, B. Engquist, and S. Chakravarthy, “Uniformly high-order accurate non-oscillatory schemes III,” *J. Comput. Phys.*, vol. 71, pp. 231–303, 1987.
- [12] C-W. Shu, and S. Osher, “Efficient implementation of essentially non-oscillatory shock capturing schemes,” *J. Comput. Phys.*, vol. 77, pp. 439–471, 1988.
- [13] G-S. Jiang, and C-W Shu, “Efficient Implementation of Weighted ENO Schemes,” *J. Comput. Phys.*, vol. 126, pp. 202–228, 1996.
- [14] G. A. Gerolymos, D. Sénéchal, and I. Vallet, “Very-high-order WENO schemes,” *J. Comput. Phys.*, vol. 228, pp. 8481–8524, 2009.
- [15] G. P. Silveira, and L. C. Barros, “Numerical methods integrated with fuzzy logic and stochastic method for solving PDEs: An application to dengue,” *Fuzzy Sets Syst.*, vol. 225, pp. 39–57, 2013.
- [16] J. W. Thomas, *Numerical Partial Differential Equations: Finite Difference Methods*. New York: Springer, 1995.
- [17] M. Motamed, C. B. Macdonald, and S. J. Ruuth, “On the linear stability of the fifth-order WENO discretization,” *J. Sci. Comput.*, vol. 47, pp. 127–149, 2011.

AUTHORS



Raphael Garcia

Raphael de Oliveira Garcia is graduated in Physics at São Paulo State University – Campus Bauru – Unesp, Brazil, Master and Ph.D. in Applied Mathematics at Institute of Mathematics, Statistics and Scientific Computing of the State University of Campinas – IMECC / Unicamp. Has experience in Applied Mathematics, Physics and Scientific Computing, acting on the following subjects: Mathematical Modeling, Numerical Analysis, General and Numerical Relativity, and Ordinary and Partial Differential Equations. Currently, he is Professor and Researcher in Department of Actuarial Science at Federal University of São Paulo – Campus Osasco – DCA / UNIFESP, accredited as professor in the Professional Master's Degree in Corporate Governance, and is one of the leaders of the Research Group, namely Laboratory for the Study of Social Risks. Furthermore, he is a participant in Nucleus for Studies in Applied Mathematics and Scientific Computing, registered with the National Council for Scientific and Technological Development (CNPq) of the Brazil.



Graciele Silveira

Graciele Paraguaia Silveira is graduated in Mathematics at Federal University of Goiás – Campus Jataí (UFG), Brazil, Master and Ph.D. in Applied Mathematics at Institute of Mathematics, Statistics and Scientific Computing of the State University of Campinas – IMECC/UNICAMP at São Paulo. Has experience in Mathematics, Applied Mathematics, Biomathematics and Scientific Computing, acting on the following subjects: Mathematical Modeling, Numerical Methods, Fuzzy Logic, and Ordinary and Partial Differential Equations. Currently, she is Professor and Researcher in Department of Physics, Chemistry and Mathematics at Federal University of São Carlos – Campus Sorocaba – DFQM/UFSCar, accredited as professor in the following postgraduate programs: Professional Master's Degree in Exact Sciences Teaching and Master's Degree in Mathematics in National Network. Furthermore, she is leader of the search group named Nucleus for Studies in Applied Mathematics and Scientific Computing, registered with the National Council for Scientific and Technological Development (CNPq) of the Brazil.

Single Phase Variable Reluctance Motor Design Using Particle Swarm Optimization

ARTICLE HISTORY

Received 16 March 2023
Accepted 12 May 2023

Danyelle Schumanski

Federal University of Paraná
Curitiba-PR, Brazil
schumanski@ufpr.br
ORCID: 0009-0008-3165-4150

Dullian Carly de Oliveira Macedo

University of Paraná
Pontal do Paraná-PR, Brazil
dullianmacedo@ufpr.br
ORCID: 0009-0004-5813-616X

Juliana Almansa Malagoli

Federal University of Paraná
Curitiba-PR, Brazil
juliana.malagoli@ufpr.br
ORCID: 0000-0002-8723-033X

Cinthia Schimith Silva Coelho

Pontifical Catholic University of Paraná
Curitiba-PR, Brazil
cinthia.coelho@pucpr.br

Single-Phase Variable Reluctance Motor Design using Particle Swarm Optimization

Danyelle Schumanski
Electrical Engineering Post-Graduation
Program
Federal University of Paraná
Curitiba-PR, Brazil
schumanski@ufpr.br
ORCID: 0009-0008-3165-4150

Dullian Carly de Oliveira Macedo
Civil Engineering Undergraduate
University of Paraná
Pontal do Paraná-PR, Brazil
dullianmacedo@ufpr.br
ORCID: 0009-0004-5813-616X

Juliana Almansa Malagoli
Electrical Engineering Post-Graduation
Program
Federal University of Paraná
Curitiba-PR, Brazil
juliana.malagoli@ufpr.br
ORCID: 0000-0002-8723-033X

Cinthia Schimith Silva Coelho
Electrical Engineering Department
Pontifical Catholic University of Paraná
Curitiba-PR, Brazil
cinthia.coelho@pucpr.br

Abstract— Electrical engines are built under electromagnetism concepts to create mechanical power, those can be seen as simple machines, as it depends on reluctance, even being called as “reluctance motor”, what gives this engine the possibility of being widely used for many purposes. The main objective of this research is to minimize copper losses of a single-phase 6x6 variable reluctance synchronous motor. For that, a Particle Swarm Optimization (PSO) algorithm will be used to obtain the optimum configuration through the Finite Elements Method (FEM). In this context, electric motor design equations were dimensioned based on similar machines. The next procedure was to use FEMM (Finite Element Method Magnetics) software, that allows the magnetic flow density analysis inside the motor air gap. Finally, it is noteworthy that the copper losses results were analyzed before and after the variable reluctance motor optimization with computational tools.

Keywords—single-phase reluctance motor, finite element method, particle swarm optimization

I. INTRODUCTION

Variable reluctance synchronous motors (VRSMs) are common engines with simple building aspects low computational and financial costs, but with a wide range of applications. These are the main reasons why reference [1] considered these motors competitive. The VRSMs dual capacity of acting like a motor and a generator, which avoids high inrush currents, reduce costs and operates under constant rotation, that are aspects of engines without windings and magnets and with a single source of incitation applied to the stator windings, what intends to minimize all the resistive losses of machinery winding that occur in the stator current flow.

Therefore, it is possible to estimate an optimal motor, based on specifications of power and size of a previous model [2]. In this context, the research objective is to minimize the copper losses, of a variable reluctance motor (VRM), by using the particle swarm method (PSO algorithm). This way, dimension parameters were used to design the original motor and the optimum motor was designed after applying the finite elements method. Finally, after the optimization process, copper losses and magnetic flow densities were evaluated and the results of both original and optimum motors were compared

II. LITERATURE REVIEW

This section will be divided in three parts: Particle Swarm Optimization, Finite Elements Method and Variable Reluctance Motor Design.

A. Swarm Particle Optimization

The PSO algorithm is a computational method that aims to optimize a problem through a proceeding that generates a sequence of approximate solutions that at each interaction, tend to converge to an optimum solution [3]. Comparative studies about the PSO (Particle Swarm Optimization) aspects were inspired on animal behavior, where each individual, of an equally dispersed population inside the problem area, efficiencies were compared. At each executed interaction, the individuals tend to group in smaller spaces around the best solution, therefore, an inertia constant can be estimated, based on the group tendency to find the best solution and considering each individual's best solution.

After a sequence of interactions performed by each individual, the optimum solution for the problem is found [4]. Next the algorithm implementation steps:

- Initially, each swarm particle has a position x_i inside the search space and the speed y_i , where the positions are automatically generated;
- For each swarm particle executed interaction, the particle position $pbest$ and the group position $gbest$ are updated, in case they are better than the previous.

During the interaction, the speeds and positions of each swarm particle are updated by the Equations 1 and 2, respectively, to obtain the new position: x_i ;

$$v(k+1) = \omega v_k + c_1 r_1 (pbest_k - x_k) + \dots + c_2 r_2 (gbest - x_k) \quad (1)$$

$$x(k+1)_i = x_k + v(k+1)_i \quad (2)$$

where $v(k+1)_i$ is the updated term for particle speed, $k+1$ is the actual instant, k is the previous instant, ω is the algorithm inertia constant, C_1 is the individual acceleration coefficient responsible for controlling the particle movement distance in one interaction, C_2 is the group acceleration

coefficient, responsible for controlling a particle movement distance in one interaction, $pbest$ is the best particle position, $gbest$ is the most visited position by the particles, r_1 and r_2 are aleatory numbers inside the research space $[0,1]$, $x(k+1)_i$ is the new particle position, x_k is the previous particle position. So, if the stop criterion is not verified, the algorithm returns to Step 2, to start the next interaction, otherwise, $gbest$ will be used as the intended solution.

It is possible to understand that the particle swarm optimization has 3 (three) main steps, knowing the number of interactions is the factor that determines the number of times where the system was evaluated and updated. The main steps, for that, are the following:

- Start: set the parameters and the population (initialize x_i and y_i randomly for each particle/individual);
- Evaluate: analyze the actual position of each particle: $gbest$ and $pbest$;
- Update: update speed and position of each particle (x_i , $t+1$ and v_i , $k+1$).

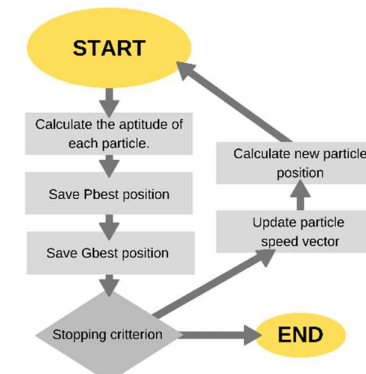


Fig. 1. PSO algorithm flowchart

Figure 1 represents the flowchart with the description of the adopted method. This way, it can be considered that the particle swarm optimization comes partially from concepts and mathematical operators for implementation. During the computational process becomes fast and of low cost in terms of speed and memory.

B. Finite Elements Method

The Finite Elements Method (FEM) is a set of different numerical methods that approximate problem solutions with Partial or Ordinary Differential Equations (PDE and ODE, respectively) through a geometry subdivision in smaller elements, known as finite elements. Knowing that the exact solution for this problem is very similar to the result of the approximate solution found through this method [5], [6], [7].

The magnetostatic design equations for magnetic circuits operation are based on Ampère and Gauss generalized laws, and they can be described as:

$$rot H = J + \frac{\partial D}{\partial t} \quad (3)$$

$$div B = 0 \quad (4)$$

Where: H is the magnetic field (A/m), J is the current superficial density (A/m²), D is the electrical induction (C/m²), and B is the magnetic flow density (T).

A magnetic field can be formed by a current in conductive materials, and due to the local form of the building relations, it is known that the flow and the magnetic field H are related, because the final magnetic permeability μ product with the field is, by definition, the magnetic flow density B :

$$B = \mu * H \quad (5)$$

The magnetic permeability describes the degree of opposition of a material to the passage of a flow [6]. In general, the relative permeability of a substance, related to the air permeability, is given by:

$$\mu_r = \frac{\mu}{\mu_0} \quad (6)$$

C. Single-Phase Variable Reluctance Synchronous Motor Design

Design equations are an essential part of the optimization processes, once the motor size affects the current flow and its distribution along the engine parts. Figure 2 shows the main engine dimensions that must be determined in order to obtain the rotor and stator projects.

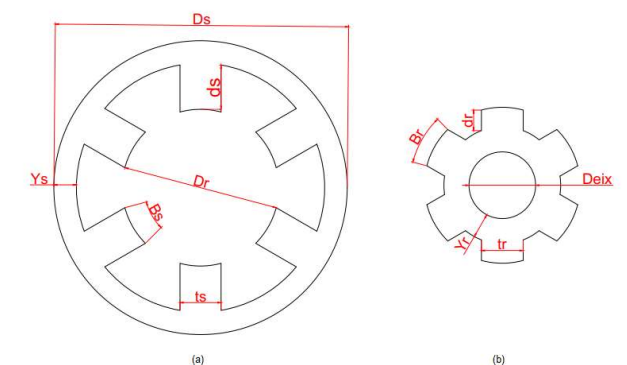


Fig. 2. Dimensions: (a) stator e (b) rotor, of the single-phase variable reluctance 6x6 motor

The rotor diameter can be calculated through the input power.

$$Dr = \left(\frac{4 * P}{TRV * k1 * w * \pi} \right)^{\frac{1}{3}} \quad (7)$$

where P is the input power; $k1$ is a format constant; TRV is the engine volume set; w is a conversion unit from RPM to rad/s [8]. And then, the pile length (L) can be considered a multiple of the rotor diameter (Dr):

$$L = k1 * Dr \quad (8)$$

Later, the number of poles inside the stator (Nps), will define the poles angles (Bs):

$$Bs = \frac{180^\circ}{Nps} \quad (9)$$

An important consideration is that the smaller the engine air gap, the highest the current flow. With this information, reference [9] considers that the engine air gap size must be close to 0,50% of the rotor diameter. Another possible consideration is to work under the percentage of 0,25%:

$$g = 0,0026 * Dr \quad (10)$$

The engine breech can be considered as the difference between the rotor diameter and the poles, and it can stand half the current that goes through the poles. Its size is defined by:

$$Ys = 1,1 * \left[\left(g + \frac{Dr}{2} \right) * \left(\sin \left(\frac{Bs}{2} \right) \right) \right] \quad (11)$$

Relations between the rotor diameter (Dr) and the external stator diameter (Ds) must be inside an interval of 0.40 to 0.70. The following relation adopted a standard of 0.55 for the relation previously defined [1]:

$$Ds = \frac{Dr}{0,55} \quad (12)$$

The poles width (ts) are another important information, that happens because the flow passes intensely through this material. This measure can be calculated using:

$$ts = 2 * \left[\left(g + \frac{Dr}{2} \right) * \left(\sin \left(\frac{Bs}{2} \right) \right) \right] \quad (13)$$

By adding the stator pole width to twice the size of the engine air gap length, the rotor pole width can be found, as can be seen:

$$tr = ts + (2 * g) \quad (14)$$

Besides that, the rotor polar arc (Br) is given by:

$$Br = 2\sin^{-1} \left(\frac{tr}{Dr} \right) \quad (15)$$

A crucial information to the motor proportions is the rotor pole height (dr), that can be obtained by:

$$dr = \frac{ts}{2} \quad (16)$$

And another important measure is the engine breech height (Ys), that can be calculated by using the width of the rotor pole (tr) with an addition of 20 to 40%. In this case, the chosen addition was of 20% [9].

$$Yr = 1,2 * \frac{tr}{2} \quad (17)$$

After calculating the rotor diameter, its height and breech, the axis diameter ($Deix$) can be calculated:

$$Deix = Dr - 2 * (dr + Yr) \quad (18)$$

And the stator pole height (ds) can be found with:

$$ds = \frac{Ds - Dr}{2} - g - Ys \quad (19)$$

However, the number of turns, per phase (Ne), is a result of the interaction between the saturated magnetic flow density ($Bsat$), the peak current (Ip) and the engine air gap width, as can be seen in the equation:

$$Ne = \frac{2 * g * Bsat}{Ip * \mu_b} \quad (20)$$

Given that the wire that must be used to form the turns has its transversal section (ac) calculated after the peak current divided by the maximum current flow (Jc) and the square root of the number of phases (q):

$$ac = \frac{Ip}{Jc\sqrt{q}} \quad (21)$$

This calculus can provide the copper loss results. The first step to understand the losses is to calculate the coil resistance:

$$Rf = \frac{\rho_0 * 2 * L * Ne}{ac} \quad (22)$$

where ρ_0 is the copper resistivity and ac the wire transversal section, what allows finding the losses of copper ($Pcopper$) and total ($Ptotal$). The first one is calculated by multiplying the coil resistivity (Rf) by any current (I_m) that passes through the motor. The total losses are copper losses in all wire turns.

$$Pcopper = Rf * I \quad (23)$$

$$Ptotal = Pcopper * Nps \quad (24)$$

At the resistive losses (or copper losses), as in the total losses, the results are estimates that can work as a good comparison between analytical and simulated values.

III. METHODOLOGY

A literature review over mathematical models for electrical motors design will be used for modeling the single-phase VRMs. With these models, a PSO algorithm will be executed under the assumptions of the finite element method. The main aim with this procedure is to reduce the copper loss inside this equipment.

Therefore, the simulations will happen through a software that develops numerical models by using electromagnetism concepts. For better understanding about purposes of this text, some items must be highlighted:

1. Set the main parameters for the motor design;
2. Define the necessary equations to model the motor;

3. Study and run the particle swarm optimization algorithm;
4. Write the objective function and the vector parameters to generate the copper loss minimization;
5. Use FEMM software to model, discretize and generate the post processing from before and after the optimization;
6. Verify and compare results between the traditional methodology of electromagnetic device design and the data generated after minimizing the objective function;
7. Generate the motor magnetic flow density and verify the copper loss before and after using the PSO.

In this context, the intention is to write the necessary equations to minimize the single-phase variable reluctance motor copper loss and, later, evaluate through FEMM software the magnetic flow density in order to achieve the research objectives.

IV. DEVELOPMENT AND DISCUSSION

In this section, discussions and results will be presented in three parts: original motor, optimization and optimum motor.

A. Original Motor

To start comprehending the problem, a single-phase motor with 6 poles was chosen, with the specifications listed on Table I. This is a model by WEG that has characteristics of an induction motor.

TABLE I. ORIGINAL MOTOR'S VARIABLES

Parameter	Values	Parameter	Values
Constant k1	1.4363	Efficiency	0.8
Set per engine volume TRV	16000	Voltage	220V
Power CV	7.5	Magnetic flow density before saturation	1.5T
RPM	1760	Magnetic flow density for saturation point	2.2T
Number of stator poles	6	Magnetic permeability	$4\pi * 10^{-7} H/m$
Input power	5500W	Current density	6 A/m ²

Therefore, a MATLAB algorithm was developed to generate the dimensions of the 6 poles variable reluctance motor stator and rotor, as well as its coils, which took to the construction of the motor via AUTOCAD with a few adjusts, and later exporting this model to FEMM, setting materials and borders. In this way, just enter the file menu and click import .dxf finally open the CAD Drawing file and the drawing will be imported into FEMM. After developing the motor project equations and running the algorithm, the main parameters were found, as shown in Table II.

TABLE II. DIMENSIONS OF THE ORIGINAL MOTOR

Parameter	Values	Parameter	Values
Dr	118.4 mm	Tr	31.34mm
L	170mm	Br	31°
Bs	30°	Dr	15.39mm

g	0.3077mm	Yr	18.8mm
Ys	16.9mm	Ds	33.16mm
Ds	215.2mm	Ne	20
ts	30.78mm	ac	5.133mm ²

In order to verify the electromagnetic device, the peak current of 53.34 A is considered, this current was obtained by the sum of the nominal current with the efficient nominal current. This means that the magnetic flow density applied to the peak current (Ip) must not have a module superior to the material magnetic flow saturation density.

The steel saturation is, approximately, of 2.2T. Figures 3 and 4 represent the motor magnetic flow density and the highest value is close to 2.2T, although inferior. In general, single-phase VRMs work close to saturation, if compared to other motors, besides that, the motor performance is still good.

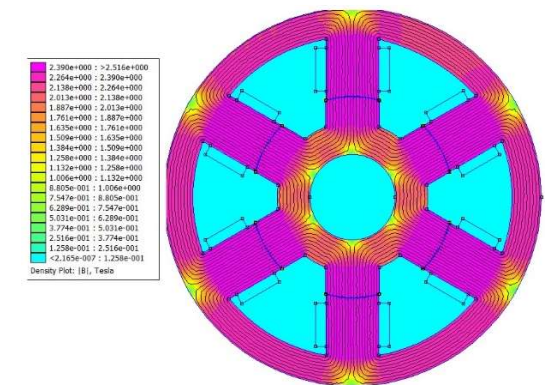


Fig. 3. Original motor aligned

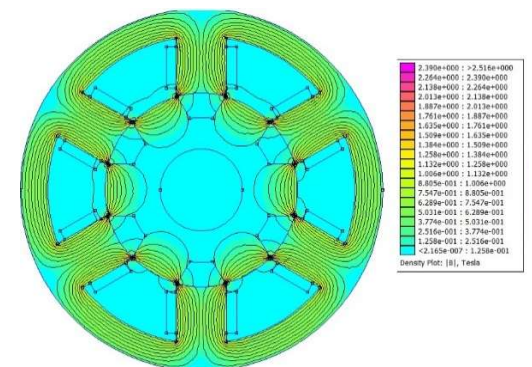


Fig. 4. Original motor misaligned

After computational simulation through FEMM software, the total copper loss was found for a peak current of 190.54W, with the analytical solution, the result is 194.52W. Therefore, the error between methods was of 2.04 %. It is concluded that with the increase of input current, the magnetic flow density is increased inside the motor air gap with a proportional relation.

The total losses had a significant scale value, which means specific optimizations will be required to reduce the power loss in the windings.

B. Optimization

Following the PSO steps and considering the Equation 24 as the objective function (OF) of the problem, the objective becomes minimizing the single-phase variable reluctance motor copper losses. Some of the variables, parameters and limits were listed on Table III.

TABLE III. VARIABLE, PARAMETER AND LIMITS

Variable	Parameter	Minimum limit	Maximum limit
x(1)	Dr(mm)	110	130
x(2)	Ys(mm)	15	20
x(3)	Yr(mm)	17	22
x(4)	Deix(mm)	30	60
x(5)	g(mm)	0.3	0.6

The Objective Function will be Equation 24, in addition, we used lateral constraints that delimit a range of variation for each design variable. It is worth mentioning some settings of the PSO algorithm, such as:

- Population size: 50;
- Generations number: 50;
- Inertia weight: 1;
- Inertia weight damping ratio: 0.99;
- Personal learning coefficient: 1.5;
- Global learning: 2.0.

Thus, the processing time was less than five minutes and the computational cost was also low, followed the settings of the notebook used in the optimization: i5, 10GB and SSD 256GB.

TABLE IV. PARAMETER VALUES AFTER THE PSO ALGORITHM IS RUN

Parameter	Optimal value
x(1)	114 (mm)
x(2)	18.4 (mm)
x(3)	20 (mm)
x(4)	45.0 (mm)
x(5)	0.305 (mm)
OF	45.2 (W)

After executing the algorithm, other two parameters were described, as well as the objective function, as seen on Table IV, where a saturation current of 27.65A and the copper losses were considered as 58.59W.

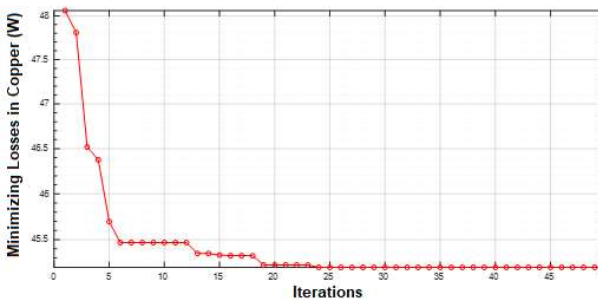


Fig. 5. Minimization of Losses in VRM Copper through PSO

In general lines, the problem is based on a mono-objective function, that with the optimization, reduced the copper loss in 29.62% when using the saturation current, as can be seen in Figure 5.

C. Optimum Motor

After running the algorithm 10 times in the test, it was possible to observe the standard deviation and mean of the design variables and the objective function. Table V shows the details of the executions, showing the optimum motor parameters.

TABLE V. PARAMETER VALUES AFTER THE PSO ALGORITHM IS RUN

Tests	OF	x(1)	x(2)	x(3)	x(4)	x(5)
1	45.0	113.8	18.3	20	45,2	0,305
2	45.2	114.0	18.4	20	44,9	0,305
3	45.3	114.2	18.3	20,1	44,9	0,305
4	45.6	114.0	18.3	20	45	0,305
5	45.3	114.0	18.4	19,9	45,2	0,304
6	45.2	114.1	18.4	19,8	45,3	0,304
7	45.3	114.0	18.3	20,1	45,1	0,306
8	45.2	114	18,2	20	44,8	0,305
9	45.1	114,2	18,5	20	44,9	0,306
10	45.0	114,1	18,4	20	45,1	0,305
Standard Deviation	0.175	0,117	0,085	0,087	0,164	0,305
Average	45.2	114,0	18,4	20,0	45,0	0,305

An interesting observation is that the measures were reduced, except for the parameters (Yr) and (Ys), the rotor and stator breeches had a measurement increase.

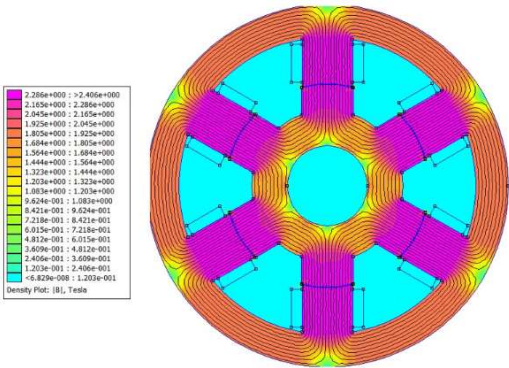


Fig. 6. Optimum motor aligned

Figures 6 and 7 indicate the magnetic flow density inside the VRM with the peak current after the optimization process.

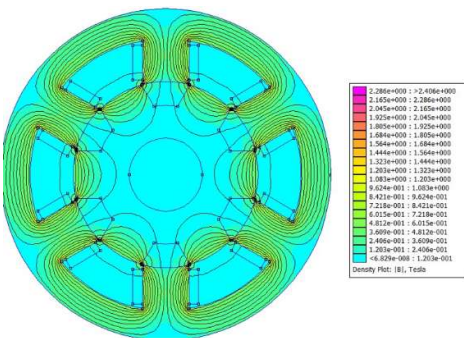


Fig. 7. Optimum motor misaligned

After the optimum motor computational simulation with FEMM software, a total copper loss for a peak current of 183.63W was found. With analytic calculus, the value was 180.78W. This means that the error between analytical and simulated results for the optimum motor is of 1.57%. It can be concluded that the magnetic flow density on the motor after the optimization reduced, what means they are under the material saturation of 2.2T. In this context, it can be observed that the PSO algorithm is effective for the convergence of the mono-objective functions, what means that the total copper loss was successfully reduced.

D. Magnetic Flow Density Analysis

The single-phase variable reluctance motors, in general, work close to the saturation, when compared to other motors, even so, their performance is good. Figures 8 and 9 represent the total airgap length when submitted to the peak current, in the original and the optimized motor.

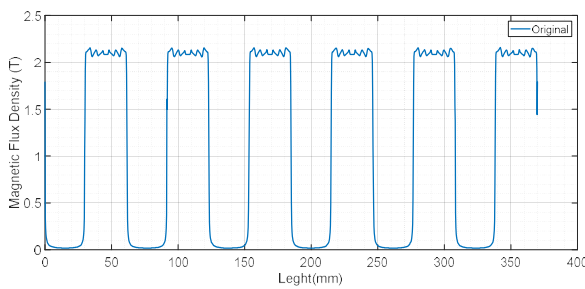


Fig. 8. Magnetic flux density for peak current of original motor

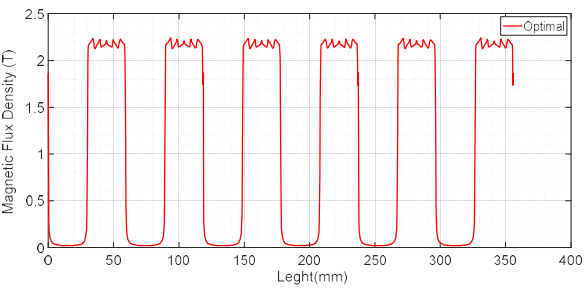


Fig. 9. Magnetic flux density for peak current of optimal motor

As the magnetic flow density in the airgap to a motor on its original and optimal unalignment, the flow density is low, therefore the data was not proven through graphics, but the flow gets thinner on the gap between stator and rotor. It can be concluded that with the increase of the input current, there is an increase to the magnetic flow inside the airgap of the single phase VRM, therefore, they have a directly proportional relation. The same behavior can be observed with the copper losses. In general, simulations stay under the iron saturation value of 2.2T, when simulated until the peak current. During the optimization procedure, the copper loss was reduced and the magnetic flow density in the airgap was under the saturation point.

V. CONCLUSIONS

Reducing the copper loss and reducing magnetic flow density, guarantees efficiency improvement during the motor

operation, in a way that the equipment can convert energy more efficiently. Comparing the properties between both motors: the original and the optimum, it can be considered that the optimization was successful, once the main objective of reducing copper loss was achieved, as well as the ideal measures were presented. It might be necessary to take into consideration the optimization of other parameters, since not all the motor measures were reduced, but some even had some increase. This is an important factor for further research.

ACKNOWLEDGMENT

The authors thank the Institutional Program of Scientific Initiation Scholarships of the Federal University of Paraná and the National Council for Scientific and Technological Development (CNPq) for their support for the development of this research and Coordination for the Improvement of Higher Education Personnel (CAPES).

REFERENCES

[1] A. E. Fitzgerald, C. Kingsley Jr. y S. Umans, Máquinas elétricas, 7th ed., Porto Alegre, Brazil: Bookman, 2014.

[2] G. Machado, «Projeto de motor a relutância variável e simulação utilizando o método dos elementos finitos (Electrical Engineering Monography),» Federal University of Uberlândia, Uberlândia-MG, Brazil, 2020.

[3] J. Kennedy y R. Eberhart, Swarm Intelligence, San Francisco, US: Morgan Kaufmann Publishers, 2001.

[4] F. Lobato, V. Steffen Jr. y A. Silva Neto, Técnicas de Inteligência Computacional com Aplicações em Problemas Inversos de Engenharia, Curitiba-PR, Brazil: Omnipax, 2014.

[5] P. Dular, «Modélisation du champ magnétique et des courants induits dans des systems tridimensionnels non linéaires (Doctorate thesis),» Université de Liège, Belgium, 1996.

[6] J. Malagoli, «Otimização multiobjetivo aplicada aos motores de indução validada via elementos finitos (Doctorate thesis),» Federal University of Uberlândia, Uberlândia-MG, Brazil, 2016.

[7] M. Luz, «Desenvolvimento de um software para cálculo de campos eletromagnéticos 3D utilizando elementos de aresta, levando em conta o movimento e o circuito de alimentação (Doctorate thesis),» Federal University of Santa Catarina, Florianópolis-SC, Brazil, 2003.

[8] A. Candido, E. Chiarello, J. Malagoli, D. Sanches y T. Terrana, «Projeto de um motor de relutância variável monofásico usando o método dos elementos finitos,» XXIII Encontro Nacional de Modelagem Computacional e XI Encontro de Ciências e Tecnologia de Materiais, XXIII ENMC and XI ECTM anais, vol. 1, pp. 1202-1211, 2020.

[9] R. Dias, «Motores a Relutância Variável 6x4 e 6x6. Estudo Comparativo de Operação e Desempenho,» Dissertation (Masters in Electrical Engineering),» Federal University of Uberlândia, Uberlândia-MG, Brazil, 2011.

AUTHORS



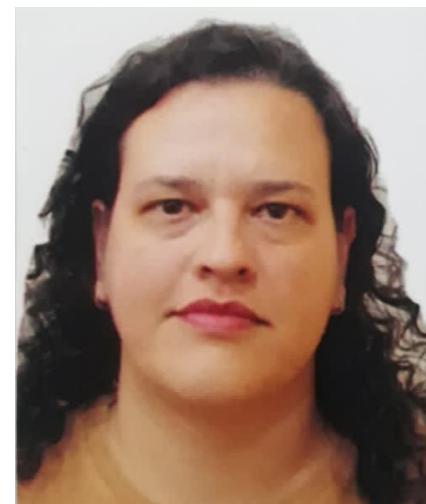
Danyelle Schumanski

Danyelle Schumanski is a civil engineer graduated by the Federal University of Paraná at the Campus Pontal do Paraná-Center for Marine Studies (CPP-CEM UFPR) and is master's degree student at the Graduate Program in Electrical Engineering at the same institution (PPGEE-UFPR). During her academic life, she has published articles in the fields of electrical engineering with magnetic field optimizations and materials science with polymeric mortars and quality analysis evaluations.



Dullian Carly Macedo

Dullian is a Civil Engineering student at the Federal University of Paraná (UFPR), who currently performs scientific initiation (PIBIC-CNPq) under the topic of PSO algorithm optimization of a 6x6 monophasic motor's airgap copper loss and magnetic flow. He's a student representative at the Civil Engineering course's collegiate at UFPR and has got a scholarship at the geoprocessing data laboratory (LAGEAMB-UFPR). Dullian is also a technician in eletromechanics by the Federal Institute of Santa Catarina (IFSC-SC) with emphasis in procedures and projects.



Juliana Malagoli

Juliana Almansa Malagoli was born in Uberlândia (Minas Gerais) and is graduated in Electrical Engineering with emphasis in computing by the Federal University of Uberlândia (UFU). Became a master in Electrical Engineering by UFU in 2012, same year when she started her doctorate thesis on the field of energy systems at the same university. In 2016 she has accomplished her doctorate degree in Electrical Engineering. Juliana has participated of the Electromagnetic Device Design and Analysis Group of the Federal University of Santa Catarina (GRUCAD-UFSC) during the period of 2012-13 and is now the lead researcher of the Study Group on Modeling, Environmental and Intelligent Systems of the Federal University of Paraná (GEMSAI-UFPR), where she is a professor of higher education at the Pontal do Paraná Campus - Center for Marine Studies (CPP-CEM UFPR) and a collaborative professor at the Graduate Program in Electrical Engineering (PPGEE-UFPR) in Curitiba.



Cinthia Silva Coelho

Has a bachelor's degree in Industrial Electrical Engineering from the Universidade Tecnológica Federal do Paraná (2013) and a master's degree in Electrical Engineering from the Universidade Federal do Paraná (2017). She is currently a professor in the Department of Electrical Engineering at the Pontifical Catholic University of Paraná. She has experience in the field of Electrical Engineering, with emphasis on Electrical and Industrial Installations, working on electrical projects for hydroelectric power plants, dams, airports, and other industrial and infrastructure facilities. Her main areas of expertise include electrical installations, lightning protection systems (SPDA), power flow, and state estimator.

Attack Taxonomy Methodology Applied to Web Services

ARTICLE HISTORY

Received 05 February 2023
Accepted 17 May 2023

Marcelo I. P. Salas
UNICAMP, University of Campinas
Campinas, SP, Brazil
marcelopalma@ic.unicamp.br
ORCID: 0000-0001-6821-0002

Attack Taxonomy Methodology Applied to Web Services

Marcelo I. P. Salas
UNICAMP, University of Campinas
Campinas, SP, Brazil
marcelopalma@ic.unicamp.br
ORCID: 0000-0001-6821-0002

Abstract— With the rapid evolution of attack techniques and attacker targets, companies and researchers question the applicability and effectiveness of security taxonomies. Although the attack taxonomies allow us to propose a classification scheme, they are easily rendered useless by the generation of new attacks. Web services, owing to their distributed and open nature, present novel security challenges. The purpose of this study is to apply a methodology for categorizing and updating attacks prior to the continuous creation and evolution of new attack schemes on web services. Also, in this research, we collected thirty-three (33) types of attacks classified into five (5) categories, such as brute force, spoofing, flooding, denial-of-services, and injection attacks, in order to obtain the state of the art of vulnerabilities against web services. Finally, the attack taxonomy is applied to a web service, modeling through attack trees. The use of this methodology allows us to prevent future attacks applied to many technologies, not only web services.

Keywords— Attack taxonomy methodology, web services, brute force, spoofing, flooding, denial-of-services, injection

Resumen— Con la rápida evolución de las técnicas de ataque y los objetivos de los atacantes, las empresas y los investigadores cuestionan la aplicabilidad y eficacia de las taxonomías de seguridad. Si bien las taxonomías de ataque nos permiten proponer un esquema de clasificación, son fácilmente inutilizadas por la generación de nuevos ataques. Los servicios web, debido a su naturaleza distribuida y abierta, presentan nuevos desafíos de seguridad. El propósito de este estudio es aplicar una metodología para categorizar y actualizar ataques previos a la continua creación y evolución de nuevos esquemas de ataque a servicios web. Asimismo, en esta investigación recolectamos treinta y tres (33) tipos de ataques clasificados en cinco (5) categorías, tales como fuerza bruta, suplantación de identidad, inundación, denegación de servicios y ataques de inyección, con el fin de obtener el estado del arte de las vulnerabilidades contra servicios web. Finalmente, se aplica la taxonomía de ataque a un servicio web, modelado a través de árboles de ataque. El uso de esta metodología nos permite prevenir futuros ataques aplicados a muchas tecnologías, no solo a servicios web.

Palabras clave— Metodología de taxonomía de ataque, servicios web, fuerza bruta, suplantación de identidad, inundación, denegación de servicios, inyección

I. INTRODUCTION

Taxonomy is the practice and science of categorization and classification of something [1]. To understand the relationship between attacks and defenses in attack taxonomy, the present research proposes a methodology to classify attacks, vulnerabilities, and faults in relation to their features.

The taxonomy proposed is constructed to build a better understanding of each attack and present possible countermeasures applied to a technology in constant development, such as web services.

Web Services are modular software applications that can be described, published, located, and invoked over a network such as the World Wide Web [2]. They are also more susceptible to security risks due to their distributed and open nature, although they provide greater connectivity, flexibility, and interoperability as one of the main benefits of this technology.

A web service [3] refers to a comprehensive set of open protocols and standards designed for seamless data exchange between various applications or systems. This versatile technology, implemented in multiple programming languages and compatible with diverse platforms, can use web services to exchange data over computer networks like the Internet in a manner like inter-process communication on a single computer. This interoperability (e.g., between Java and Python or Windows and Linux applications) is due to the use of open standards, and these features make them more attractive while generating new challenges for maintaining information security.

A study quoted in [4] described that the use of web service technology reopened 70% of the vulnerabilities filtered by firewalls. Also, in addition to traditional vulnerabilities, there are new ones in web services that must be considered. According to the OWASP¹ Top 10 [5] and the CWE² [6], injection and denial of service attacks were among the most exploited in 2021.

The Simple Object Access Protocol (SOAP), as shown in Fig. 1, used to exchange messages among participants, does not address security by itself, and it can bypass a firewall through the Extensible Markup Language (XML), usually via HTTP [7]. The receiving system interprets the message and sends back a response in the form of another SOAP message under a Service Oriented Architecture (SOA). This vulnerability allows attackers to easily exploit and manipulate the messages to their advantage.

Some research [5], [6], [8], [9] analyzed the web service attacks as a set of threats, which can be mitigated by current security specifications, such as WS-Security, XML-Encryption, XML-Signature, among others. The problem resides in having a partial view of the attack and difficulty classifying it, selecting characteristics of system security

based on a taxonomy tested to create a testing plan for web services, and using appropriate countermeasures.

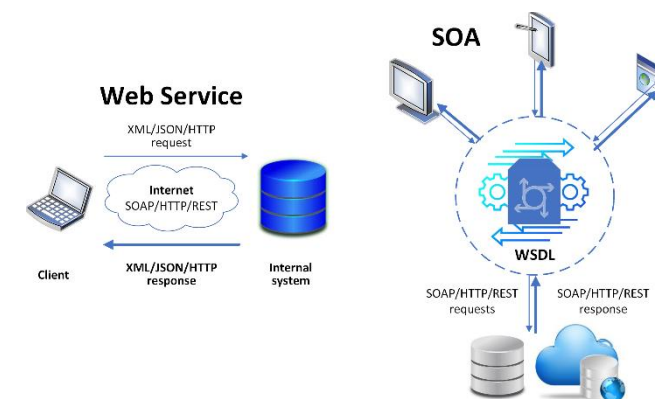


Fig. 1. Illustrates the process of client invokes a web service by sending an XML request services, which then sends back an XML response. It uses many standards such as WSDL3 and SOAP.

This proposal designs and applies a methodology to classify thirty-three (33) web services attacks into five (5) categories, describing their security properties affected, the level at which they develop, and countermeasures, among others, in the section III. For ease of understanding, attack trees are applied to model vulnerabilities against web services.

The rest of the paper is organized as follows. Section II provides a review of related work and existing attack taxonomies. Section III proposes a methodology for attack taxonomy delves into security challenges and attacks on web services with possible countermeasures. Section IV applies the methodology of attack taxonomy. Section V concludes the research, describing the contributions and future work.

II. RELATED WORK

Taxonomy [10] is described as the study of the common principle of scientific classification and includes basis, principles, procedures, and rules. It can be used to indicate the actual categorization of objects, e.g. attack taxonomy that describes vulnerabilities like injection or brute force attack. This section provides an overview of the properties of a taxonomy as well as existing works.

A. Analysis of the Properties of the Taxonomy

According to [11], it is important to define the properties and requirements for a correct classification process in an attack taxonomy, defined as follows:

- Acceptable by the security community.
- Comprehensible taxonomy could be understood by security experts.
- Completeness and exhaustiveness ensure that all types of attacks are covered, and if new ones emerge, the taxonomy can be expanded to include them.
- Determinism is the procedure by which classification occurs and is clearly defined.
- Mutually exclusive means that each attack can only be categorized, at most, into one category.

- Repeatable means that the classification of an attack should be reproducible.
- Constant and defined security terminology means that it makes use of standard and well-established nomenclature in the area.
- Well-defined terms allow a categorization of attacks through precise features.
- Unambiguous means that the taxonomy must have clearly defined classes.
- Usefulness is a requirement that can currently be tested through security testing.

B. A Brief Survey of Attack Taxonomies

Numerous attack taxonomies have been developed over the years to protect web applications and services. Below, we describe these taxonomies based on their potential and relationship with this research.

Chan et al. [14] presented a taxonomy that offers a comprehensive framework for understanding attacks within a new classification. The taxonomy organizes and classifies attacks based on three key parameters: the web services layer, attack methodology, and impact. This proposed taxonomy provides the necessary flexibility to classify emerging web service attacks within a Service-Oriented Architecture (SOA) environment.

Karumanchi and Squicciarini [7] went beyond addressing commonly known web-based vulnerabilities like SQL injection and session replay. They also conducted an examination of web service-specific vulnerabilities, highlighting the potential for attacks arising from subpar service construction and inadequate maintenance. In their comprehensive analysis, they classified each vulnerability based on a novel taxonomy, discussing potential solutions and associated impacts. Additionally, they proposed real-time analysis-based detection methods. However, it should be noted that their taxonomy does not include the essential tools for the classification of attacks.

Derbyshire et al. [1] applied two approaches to evaluate seven taxonomies. The first approach analyzed the criteria used for the taxonomy creation and critical components. The second applied historical attack data to each taxonomy under review, more specifically, attacks in which industrial control systems have been targeted. This combination of methodologies enables a comprehensive exploration of existing taxonomies, offering insights from both theoretical and practical perspectives, thus fostering a deeper understanding of the subject matter.

Panchal et al. [12] introduced an Industrial Internet of Things (IIoT) attack taxonomy as a valuable resource for mitigating attack risks. Their research incorporated four key dimensions: attack vector, attack target, attack impact, and attack consequence. These dimensions provide a comprehensive framework for modeling attacks on industrial infrastructures and offer insights into the attack methodology, affected components, and the attacker's objectives. Nevertheless, the taxonomy did not include a breakdown of attacks into specific features, nor did it facilitate reproducibility of the attacks.

¹ The Open Web Application Security Project (OWASP) is a nonprofit foundation dedicated to improving software security.

² Common Weakness Enumeration (CWE) is a universal online dictionary of weaknesses that have been found in computer software.

Simmons et al. [13] employed a tree structure in their taxonomy, which encompassed five primary categories: attack vector, operational impact, defense, informational impact, and target (AVOIDIT). The authors placed significant emphasis on the classification of attacks, primarily focusing on the attack processes. Consequently, their attention did not extend to mitigating the attacks or incorporating other crucial aspects, such as the security properties at risk (confidentiality, integrity, and availability).

Chan et al. [14] exposed a taxonomy that provides a way to understand attacks on a new classification. Attacks were grouped and classified based on three parameters: the web services layer, attack methodology, and effect. The proposed taxonomy can provide the flexibility to classify new web service attacks in a SOA environment only.

In [27] the authors focus on providing a comprehensive understanding of security attacks and risk assessment in the context of cloud computing. The authors aim to develop a taxonomy that categorizes various security attacks in the cloud environment and propose a risk assessment framework for evaluating the associated risks.

Yassine et al. [28] aim to provide an organized and structured taxonomy that categorizes various types of threats faced by web applications. The authors develop a taxonomy that classifies web application threats into different categories, considering factors such as the attack vector, the targeted component, and the impact of the threat.

Prinetto and Roascio [29] present a thorough and structured taxonomy of hardware vulnerabilities and attacks. This research addresses hardware trust and the authenticity of components, proposing a comprehensive taxonomy of hardware vulnerabilities and attacks, classifying them based on their domain, nature, target, goal, implementation method, and domain.

Previous research describes different approaches and features of taxonomies created to analyze vulnerabilities in web services and applications. This research has several benefits compared to the research cited above:

1. A specific approach to the challenges of classifying attacks in web services.
2. The proposed methodology considers the constantly evolving nature of attacks, recognizing that attackers are always generating new techniques and attack schemes.
3. The inclusion of multiple categories and attacks provides a wide range of attacks categorized based on attack properties, making them easy to classify.
4. The application of attack trees allows us to model and visualize the sequence of attacks and the interrelationships between them, providing a graphical representation and understanding of how they can be mitigated and counteracted.
5. Properly classifying attacks using this methodology helps researchers determine which countermeasures are most effective for each type of attack. By providing this information, the research makes it easier to select and apply the appropriate countermeasures to protect web services or another technology.

III. PROPOSED METHODOLOGY

One of the difficulties in finding vulnerabilities in web services during the execution phase is identifying attack scenarios. These scenarios are time-consuming to find and set up a bank of relevant attacks and automate them according to the test environment.

The objective of this research is to identify new types of attacks against web services following the steps outlined in Fig. 2. The rest of the section describes the steps of the attack taxonomy methodology.

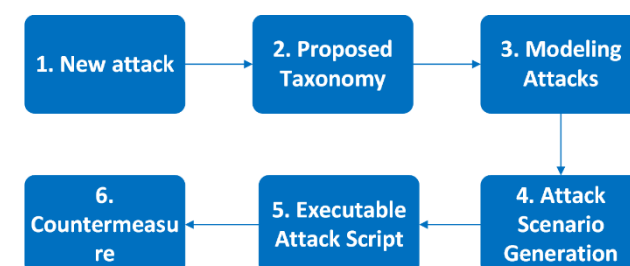


Fig. 2. Attack taxonomy methodology applied to web services.

A. New Attack

There are two ways to identify new vulnerabilities in web services: detecting existing attacks in the environment or identifying new vulnerabilities in the protocol stack used by the web service. In the first case, the researcher has to make use of honeypots and websites that publish attacks. In the second case, the researcher must use black-box or white-box techniques to identify unknown vulnerabilities.

The first step is to define an approach to systematically obtain new attacks, or variations thereof, that is successful enough to be classified in the attack taxonomy in the next step. If the attack is like another known attack, the researcher must decide whether to create a new item or a subclassification item, based on the characteristics described in Table I.

B. Attacks Taxonomy Against Web Services

Due to the magnitude of threats against web services today, it is possible to use different taxonomies to classify these attacks. This allows a better understanding of the potential threat and facilitates the application of possible countermeasures to each of these. For example, Landan in [15] categorizes security challenges as threats, attacks, and security problems divided into two levels:

- Service-level threats, also known as process-level attacks, are carried out in two web service protocol stacks (service discovery (UDDI) and service description (WSDL)) as well as SOAP message processing. Among the attacks against WSDL and UDDI are malicious code injection, phishing, denial-of-service (DoS), XML spoofing schema, session hijacking, and others.
- Message-level threats executed on the other two protocol stacks: transport protocol (HTTP, SMTP, FTP) and message protocol (XML, WS-Addressing, SOAP). These protocols enable attackers to execute various malicious activities, including fault injection attacks, message forwarding, message validation manipulation, interception, and compromising message confidentiality, among other potential exploits.

The problem with the classification above consists of the inability to clearly define which protocol stack the attack belongs to. Another way to classify the attacks comes through the security properties affected. These properties are confidentiality (C), data integrity (I), availability (A) and access control issues (AC). In this case, the vulnerability occurs when an attack violates more than one property, e.g., a WSDL Scanning attacks violate confidentiality and access control properties because it looks for vulnerabilities in the WSDL.

The design of the following taxonomy allows researchers to explore many types of vulnerabilities and use specific features of each attack to analyze how to affect web services, as shown in Fig. 3, of a tree model of attacks against web services composed of 33 attacks classified into five categories.

The selection of the number of attacks and categories in this research is justified based on the security properties affected by the attacks, references to many kinds of known attacks, level of attack (WSDL or SOAP), impact level according to the OWASP, type of attack that can concentrate various types of known attacks, and possible countermeasures according to the WS-I. It is also possible to increase both the number of attacks and the categories of attacks.

These threats were collected from several studies that examined potential vulnerabilities, even with the use of security specifications such as WS-Security or WS-Trust. The Table I describes these attacks and its features, using the following criteria:

- Attack Type: Denial-of-Services (DoS), Brute Force (BrF), Spoofing (Spo), Flooding (Flo) and Injection (Inj).
- Attack: number and name.
- References.
- Security properties affected: confidentiality (C), integrity (I), availability (A), access control (AC).
- Attack level: service level (WSDL), and message level (SOAP).
- Sending requests to execute the attack, e.g. 1 or 1+ (at least one or more messages) or n (many messages).
- According to the OWASP [9], the impact level for successful attacks in a business environment is classified as low, medium, or high risk.
- Possible countermeasures, according to the Web Services Interoperability Organization (WS-I) [16].

Below, it is described each of these types of attacks against web services, as depicted in Fig. 3.

1) *Denial-of-Service Attacks (DoS)*: It is an attempt to make the system resources unavailable to its users. This is not an invasion of the system, but an invalidation by overload. DoS attacks are typically carried out in two

ways: i) forcing the victim system to reboot or consume all resources, such as memory or processing overhead, so that it cannot provide its service; and ii) interfering with communication between users and the target system in order to impair its functioning. It is composed of replay attacks, oversized payloads, coercive parsing, oversize cryptography, attack obfuscation, XML bombs, invalidated redirects and forwards, SOAP attachment, and schema poisoning.

2) *Brute Force Attacks*: The strategy used by this attack is to break the security system's encryption, consisting of exploring all possible key combinations in a cipher algorithm until the correct key is found. These attacks, since they use the method of trial and error, are very expensive in computational time. Web services that are vulnerable to this attack are insecure cryptographic storage, broken authentication, and session management.

3) *Spoofing Attacks*: This attack consists of a set of identity theft techniques in which the hacker successfully masquerades as another in order to falsify data and thereby gain an illegitimate advantage, i.e., web service resources. It is composed of SOAPAction, WSDL scanning, insufficient transport layer protection, WS-Addressing, middleware hijacking, metadata, security misconfiguration, unauthorized access, routing detours, attacks on WS-Security, attacks on WS-Trust, and malicious content.

4) *Flooding Attacks*: It is characterized by trying to cause a breakdown in the target system by providing more workload than the system can support. A flooding attack uses traffic redirection techniques and output port modification. It is composed of instantiation flooding, indirect flooding, and BPEL state deviation.

5) *Injection Attacks*: This type of attack involves intercepting and manipulating messages. The attackers aim to exploit vulnerabilities on the server-side to execute malicious commands, gain unauthorized access to data, and take control of the server. Injection attacks encompass various subtypes, such as XML injection, SQL injection, XPath injection, cross-site scripting (XSS), cross-site request forgery (XSRF), fuzzing scans, invalid types, parameter tampering, malformed XML, and Frankenstein messages (timestamp tampering).

In the first two steps, as shown in Fig. 2, the researcher can verify whether or not the new attack falls into any of the categories and types of attacks since there is a high possibility that it will be ruled out as a known attack and move on to the countermeasures phase

TABLE I. ATTACKS AGAINST WEB SERVICES AND THEIR COUNTERMEASURE.

Type	Attack	References	Properties	Attack Level	Request	Impact	Counterme. (use)
DoS	A.01. Replay Attack	[17], [18]	A	Message	n	Low	-
DoS	A.02. Oversize Payload	[5], [18], [19]	A	Message	1	Low	-
DoS	A.03. Coercive Parsing (Recursive Payloads)	[5], [18], [19]	A	Message	1	Medium	-
DoS	A.04. Oversize Cryptography	[5], [18], [19]	A	Message	1	Medium	-
DoS	A.05. Attack Obfuscation	[5], [18], [19]	A	Message	1	Low	-
DoS	A.06. XML Bomb	[17], [18]	A	Message	1	Low	-
DoS	A.07. Invalidated Redirects and Forwards	[5], [18]	C, A	Message	1+	Medium	-
DoS	A.08. Attacks through SOAP Attachment	[17], [18], [20]	I, A, AC	Message	1	High	-
BrF	A.09. Insecure Cryptographic Storage	[5], [21]	C, I, A	Service	1+	Medium	-
BrF	A.10. Broken Authentication and Session Mgmt	[5], [21]	AC	Service	1+	High	WS-Security
Spo	A.11. WSDL Scanning	[19]	C, AC	Message	1	Medium	WS-Security
Spo	A.12. Insufficient Transport Layer Protection	[5]	C	Message	1	High	WS-Security
Spo	A.13. Metadata Spoofing	[19]	All	Message	1+	Medium	-
Spo	A.14. WS-Addressing Spoofing	[19]	C, A	Message	2	Medium	WS-Security
Spo	A.15. Middleware Hijacking	[19], [21]	A, AC	Service	2+	Medium	-
Spo	A.16. SOAPAction	[19], [20]	AC	Message	1	Medium	WS-Security
Spo	A.17. Security Misconfiguration	[5]	All	All levels	1+	Medium	WS-Security
Spo	A.18. Routing Detours	[21], [22]	C, I	Message	1+	Medium	WS-Security
Spo	A.19. ATK on WS-Trust, WS-Secure Conversat.	[20]	I, AC	Message	2+	Medium	WS-Security
Spo	A.20. Attack on WS-Security	[20]	AC	Message	2+	High	WS-Security
Flo	A.21. Instantiation Flooding	[19]	A	Service	1	High	-
Flo	A.22. Indirect Flooding	[19]	A, AC	Service	2+	High	-
Flo	A.23. BPEL State Deviation	[19]	A	Service	1	High	-
Inj	A.24 XML Injection	[5], [19]	I	Message	2+	Low	WS-Security
Inj	A.25 SQL Injection	[5], [21]	I	Message	2+	High	WS-Security
Inj	A.26 XPath Injection	[17]	I	Message	2+	High	WS-Security
Inj	A.27 Cross-site Scripting (XSS)	[5], [21]	I	Message	2+	Medium	WS-Security
Inj	A.28 Cross-site Request Forgery (XSRF)	[5], [21]	AC	Message	2+	Medium	WS-Security
Inj	A.29 Fuzzing Scan	[17]	I	Message	n	Low	WS-Security
Inj	A.30 Invalid Types	[17], [19], [21]	I	Message	2+	Low	WS-Security
Inj	A.31 Parameter Tampering	[5], [17], [21]	I	Message	2+	Low	WS-Security
Inj	A.32 Malformed XML	[17]	I	Message	2+	Medium	WS-Security
Inj	A.33 Frankenstein Message (Modify Timestamp)	[20]	I, AC	Message	1+	Medium	WS-Security

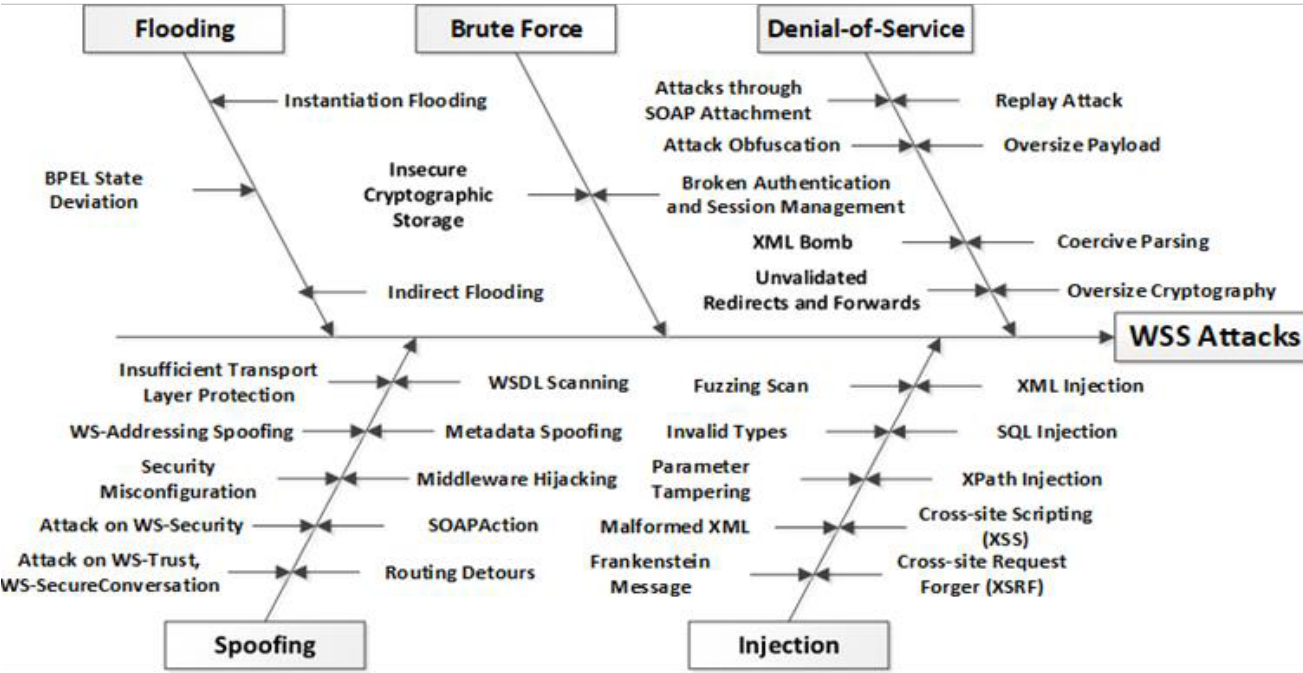


Fig. 3. Web services security attacks composed of 33 attacks classified into 5 categories

One of the challenges in reducing potential vulnerabilities in web services through detecting malicious or accidental flaws is determining the appropriate attack scenarios. At this point, it is possible to automate many types of attacks according to the test. This technique is described in the next step.

C. Modeling Attacks

To model the attacks, the researcher makes use of a modeling tool like SecureITree [23]. This tool allows building attack trees according to the attributes proposed in Subsection III-B.

Initially, the attacks are classified according to some features described in Table I. These features allow generating an attack tree, as seen in Fig. 3, to allow the researcher to select an attack type, attack its main security properties, and identify other characteristics that allow a successful attack.

A starting point is to develop a set of questions that allow the researcher to identify if they can carry out the attack on their systems. Next, it is suggested to ask the following questions:

- Does the researcher have the ability (knowledge) to carry out the attack?
- Does the researcher get to emulate the attack scenario through a tool or platform such as SoapUI [24]?
- Does the web service satisfy the required features to carry out the attack?
- Is WS-Security or another standard protecting the web service from this attack (impossible) or not (possible)?

The researcher must answer the questions for each attack. If the four questions are affirmative, i.e., possible <P, P, P, P>, the attack can be executed. There is a special case when it is necessary to use WS-Security to execute a certain type of attack, such as Oversize Cryptography or Attack Obfuscation. In this case, the third attribute will have the value of P (possible) only when the web service executes the security standard that allows the attack execution. Otherwise, it will be impossible (I). Besides, the security standard must not prevent the execution of the attack, or the fourth attribute will be impossible (I).

Once the four questions for the attack taxonomy (see Table I) have been evaluated with logical values (possible and impossible), the attack tree will be obtained and it can begin to look for vulnerabilities in web services. The output can be seen in Fig. 4.

After ensuring the requirements of the four attributes of the selected attacks, the researcher can generate the scenarios for each of them

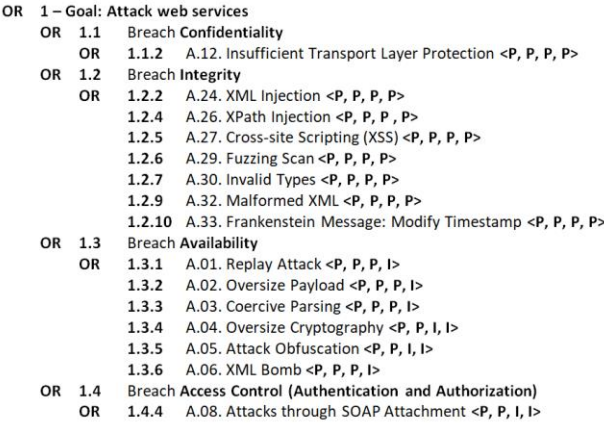


Fig. 4. Web services attack tree.

D. Attack Scenario Generation

The attack scenarios are produced automatically through the attributes used in the attack tree. The scenarios obtained in this section can be used to create a library of attacks useful for testing other web services, protocols, or systems, facilitating their reuse. Fig. 5 describes an example of an attack scenario for the XML Injection attack, using information obtained in [5], and [19] about the operation of the attack and the requirements.

The result of this step is the generation of attack scenarios described in textual language, which is on the same level of abstraction as the attack tree. This descriptive format proves valuable for test analysts and security experts due to its ease of configuration. However, it is important to note that this type of description is not directly processable by a tool or automated system.

Analysts must perform a set of refinement steps in order to transform attack scenarios in textual language into a script executable by the attacker's preferred tool, such as SoapUI [24].

1	Objective: XML Injection attack seeks to modify XML structure.
2	Preconditions:
3	- The client sends a request to Web Services via a SOAP message.
4	- The client does not use a secure communication scheme.
5	- The WSDL describes at least one parameter.
6	Attack:
7	AND 1. If required a WSDL to send a message.
8	2. If it contains the searched <String> string or tag.
9	3. Modify the request parameters by a non-default message.
10	4. Generate the new SOAP message.

Fig. 5. XML Injection features for the Attack Scenario.

E. Executable Attack Script

The generation of executable attack scripts is important to experimentally validate the vulnerability found. The messages exchanged between the server and the client must be monitored and collected to determine if the attack was effective (true positive) or not (true negative), as well as if it is necessary to modify the attack script.

It is recommended the use of soapUI [24] or another security testing tool to automate the attacks and allow to monitor in real time, as seen in the Fig. 6.


```
1: HTTP/1.1 200 OK
2: Cache-Control: private, max-age=0
3: Content-Length: 455
4: Content-Type: text/xml; charset=utf-8
5: Server: Microsoft-IIS/7.0
6: X-AspNet-Version: 4.0.30319
7: X-Powered-By: ASP.NET
8: Date: Mon, 26 Oct 2021 10:12:06 GMT
9:
10: <soap:Envelope xmlns:soap="http://schemas.xmlsoap.org/soap/envelope/"
11: xmlns:xsi="http://www.w3.org/2001/XMLSchema-instance"
12: xmlns:xsd="http://www.w3.org/2001/XMLSchema">
13:   <soap:Body>
14:     <GetISOCountryCodeByCountyNameResponse xmlns="http://www...X.NET">
15:       <GetISOCountryCodeByCountyNameResult>&lt;NewDataSet />
16:     </GetISOCountryCodeByCountyNameResult>
17:   </GetISOCountryCodeByCountyNameResponse>
18: </soap:Body>
19: </soap:Envelope>
```

Fig. 6. The vulnerability was discovered through an injection attack and the execution of the string <NewDataSet.

Finally, it is important to reduce potential false positives and false negatives through the development of rules that allow for determining when there is a confirmed attack.

F. Countermeasures

During the 2000s and 2010s, one of the most widely used options against cross-site scripting (XSS) or XML Injection attacks was the use of security protocols developed by OWASP, such as WS-Security or XML Encryption [25]. In recent years, the use of Web application firewalls (WAF) [26] became popular. This tool is a specialized application firewall designed to filter, monitor, and block HTTP traffic to and from a web service, with a particular emphasis on Layer 7 applications. It serves as an effective defense against attacks that exploit well-known vulnerabilities in web applications, including SQL injection, cross-site scripting (XSS), XPath Injection, and malformed XML. By leveraging its capabilities, this tool acts as a proactive shield, preventing the successful exploitation of known vulnerabilities of a web application.

Although WAF can block different types of attacks, it still requires the help of the authentication and encryption protocols offered by SOAP Foundation.

This section describes some security mechanisms to protect web services (see Fig. 7) against many kinds of web service attacks like WS-Security or XML Encryption taking into account three aspects: (i) message authentication in order to make sure that a transaction between the server and its client is legitimate; (ii) confidentiality to protect exchanged messages against interception by an unauthorized third party; and (iii) integrity of messages sent between server and client in order to remain unaltered.

X = Partial protection of this kind of attacks. O = Reduce the impact of this kind of attacks.						
	DoS	Brute Force	Flooding	Spoofing	Injection	XML Injection
XML Encryption			O	X	O	
XML Signature			X	X	O	X
Security Tokens			X	X	O	O
WS-Addressing					O	O
SSL/TLS [end to end]			X	X	O	X
HTTP Authentication			X	X	O	

Fig. 7. Security mechanisms to protect web services.

Some security mechanisms to protect web services are described below.

1) WS-Security (WSS): This standard contains specifications that guarantee the confidentiality and integrity of messages and user authentication. It inserts a layer over the SOAP message to build more secure and robust services with broad interoperability. In addition to being a solid and open security model, WS-Security is fast-developing, allowing users to encrypt XML documents and secure sessions between two or more parties [25] using other specifications such as Security Tokens, XML-Encryption, and XML-Signature.

2) Security Tokens: The Security Token is a security specification for providing authentication and authorization in web services, providing access rights to application servers. It makes use of the tag to provide different credential types, such as identification by user/password, to more complex ones, based on certificates such as X.509 and Kerberos [25].

3) XML Encryption (XML-Enc): This specification provides confidentiality and authentication to the web service by encrypting information between the parties. It makes use of the <EncryptedData> tag to use the encryption key. Thus, users who do not have the key will not have access to the message and its contents. The technology [25] allows the use of multiple cryptographic keys for different parts of a message. In turn, the same message can have several receivers, and each receiver only has access to its own parts of the message.

4) XML Signature (XMLDSig or XML-Sig): This pattern makes use of the <Signature> tag and is used for two purposes: it validates Security Token credentials and verifies that messages are not modified during transmission to ensure their integrity. Verification of credentials is done using the signature in combination with the certificate to ensure that the user is who they claim to be [25]. Similar to XML Encryption, XML-Sig allows users to sign certain portions of the message.

IV. ATTACK TAXONOMY METHODOLOGY APPLIED TO SESSION FIXATION

A recently published attack [30], called Session Fixation, is used to apply the attack taxonomy methodology.

A session fixation attack occurs when an attacker manipulates a user's session ID to a specific value, granting them unauthorized access. Attackers employ various techniques, such as cross-site scripting exploits or reusing HTTP requests, to achieve session fixation. The attack typically unfolds in two steps: the attacker fixes the victim's user session ID and then the victim unknowingly logs in, unknowingly exposing their online identity. With the fixed session ID value, the attacker can subsequently hijack the victim's user identity and gain unauthorized control.

In this way, the attacker modifies the message (integrity) and attempts to impersonate the identity of the victim (access control) by sending one or more messages at the SOAP message level. According to OWASP [9], the impact level of a successful attack is classified as medium to high risk for reaching a level of authentication in the system. Furthermore, it can be classified as a spoofing attack in the subsection III-B3 due to its similarity to A.17 Security Misconfiguration, described in Table I.

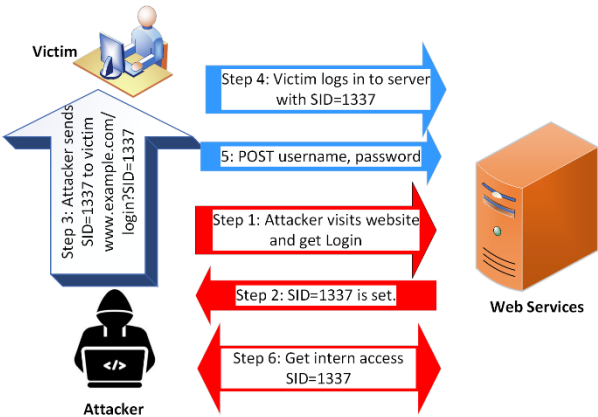


Fig. 8. Exemplified Session Attack in web services

The Session Fixation attack would fall under the category of "Spoofing" (Spo). The main goal of a Session Fixation attack is to trick the system into accepting a fake session or identity as valid. This involves intentionally manipulating or fixing a user's session ID and then impersonating that identity. This type of attack is considered a form of identity theft and falls under the category of "Spoofing".

In the modeling phase, the attacker gets <p, i, p, p> because:

- The attacker has knowledge to implement the attack (possible), as shown in Fig. 8.
- SoapUI does not emulate the Session Fixation attack (impossible). In this case we use another tool.
- I have access to web services that require authentication and do not use a web application firewall or other protection standard (possible).
- WS-Security, XML-Encryption, XML-Signature, and Security Tokens can protect against this attack (possible).

Since the attacker cannot ensure the requirements for the four attributes described above, the Session Attack must be manually programmed to generate the attack scenario. In generating the attack scenario, the attacker can use Burp Suite³, OWASP ZAP⁴, WebScarab⁵, or BeEF⁶ to simulate the type of attack. Next, we make the textual description of the Session Fixation attack described in Fig. 9.

In the next step, we generate the script for the Session Fixation attack using cookies in a web server authentication scenario:

1. Attacker: Generates a fixed session ID: "fixed_session_id".
2. Victim: The victim visits the targeted website without having previously logged in.

1	Objective: Take control of a user's session on a website.
2	Preconditions:
3	- WS must use sessions to authenticate and maintain user state.
4	- WS must use a session ID mapping mechanism, such as cookies, URLs, or hidden fields in forms.
5	- The attacker needs to know the mechanism used to assign and manage session IDs in the WS.
6	- The attacker must have the ability to manipulate or force the user to use a specific session ID before they log in to the WS.
7	Attack:
8	AND 1. The attacker identifies the mechanism used by the WS to assign and manage session IDs.
9	2. The attacker generates a specific session ID and sets it, either by manipulating the allocation mechanism or tricking the user into using a predefined session ID.
10	3. The attacker sends the pinned session ID to the legitimate user or otherwise makes it accessible for the user to use when logging into the WS.
11	4. The user logs into the WS using the session ID provided by the attacker.
12	5. Once the user is logged in, the attacker uses the same set session ID to access the user's active session and carries out malicious actions on its behalf.

Fig. 9. Session Fixation features for the Attack Scenario.

3. The victim receives the email and clicks on the malicious link. The victim's browser opens the target website and sets a session cookie with the fixed session ID provided by the attacker.
4. The attacker now knows the session ID that the victim will use when logging in.
5. The victim continues to use the website and decides to log into her account. The victim's browser sends an HTTP POST request to the server with the login credentials and the session cookie containing the fixed session ID.
6. Server (response at HTTP protocol level):
 - The server receives the POST request with the login credentials and the session cookie.
 - The server validates the credentials and verifies the session cookie.
 - Because the session cookie is valid and contains the fixed session ID, the server considers the session to be legitimate and authenticates the user as the victim.

In this way, the attacker managed to set the victim's session ID before the victim logged in. As a result, the attacker can impersonate the victim and access the victim's account without having to provide the correct login credentials.

At the HTTP protocol level, the POST request sent by the victim's browser to the server would include the login form data and headers, such as the following:

```
1 POST /login HTTP/1.1
2 Host: www.bank.com
3 Cookie: session_id=fixed_session_id
4 Content-Type: application/x-www-form-urlencoded
5 Content-Length: 23
6
7 username=victim&password=secret
```

Fig. 10. The session fixation attack was successful through an HTTP protocol request. The attacker has access to the user name and password of the victim.

³ <https://portswigger.net/burp>
⁴ <https://owasp.org/www-project-zap/>

⁵ <https://github.com/OWASP/OWASP-WebScarab>
⁶ <https://beefproject.com/>

The server would process the request and authenticate the user based on the session cookie containing the fixed session ID provided by the attacker.

There are several countermeasures that you can use to mitigate a Session Fixation attack. Since this attack undermines the properties of integrity (I) and access control (AC), Fig. 7 of security mechanisms suggests the use of a combination of techniques to reduce the impact of said attack, made up of XML Encryption, XML Signature, Security Tokens, use of SSL, and HTTPS. We can also establish a session expiration policy, implement anti-CSRF tokens, and monitor and log suspicious activities, both on the server and on the compromised computer.

V. CONCLUSIONS

The attack taxonomy methodology contributed to the development of security research in web services by describing the security properties affected, the level at which they develop, and other features.

This methodology can be used to explore many types of vulnerabilities and use specific features of each attack, like Session Attack. The objective is to analyze how an attack can affect web services, in addition to creating new attacks and selecting possible countermeasures.

In this way, this research described five categories of web services attacks (brute force, spoofing, flooding, denial-of-services, and injection attack types) along with thirty-three (33) attacks to provide a state of the art.

As shown in Table I, this taxonomy allows researchers to classify new attacks based on properties (integrity, availability, confidentiality, and access control), level of attack (WSDL or SOAP), amount of exchange of messages, or level of impact according to the OWASP Top Ten.

Furthermore, a correct classification or grouping of an attack will allow researchers to more easily determine which potential countermeasures to employ.

In the future, it is proposed to apply this systematic methodology to different technologies. Furthermore, it is possible to combine this methodology with malware attacks like botnets.

ACKNOWLEDGMENT

The author would like to thank Eliana Martins (LSD Laboratory), and Paulo Lício de Geus (LASCA Laboratory) from the Computing Institute of University of Campinas (UNICAMP).

REFERENCES

- [1] R. Derbyshire, B. Green, D. Prince, A. Mauthe, and D. Hutchison, "An analysis of cyber security attack taxonomies," in 2018 IEEE European Symposium on Security and Privacy Workshops (EuroS&PW), pp. 153–161, IEEE, 2018.
- [2] C. Ferris and J. Farrell, "What are web services?," *Communications of the ACM*, vol. 46, no. 6, p. 31, 2003.
- [3] "Web services tutorial." https://www.tutorialspoint.com/webservices/what_are_web_services.htm, 2022.
- [4] I. Siddavatham and J. Gadge, "Comprehensive test mechanism to detect attack on web services," in 2008 16th IEEE International Conference on Networks, pp. 1–6, 2008.
- [5] D. Wichers and J. Williams, "Owasp top-10 2021." <https://owasp.org/Top10/>.

- [6] C. Top and M. Dangerous, "2021 cwe top 25 most dangerous software weaknesses, 2021." https://cwe.mitre.org/top25/archive/2021/2021_cwe_top25.html.
- [7] H. Yuan, L. Zheng, L. Dong, X. Peng, Y. Zhuang, and G. Deng, "Research and implementation of web application firewall based on feature matching," in *International Conference on Application of Intelligent Systems in Multi-modal Information Analytics*, pp. 1223–1231, Springer, 2019.
- [8] B. Jagruti, P. Nidhi, and D. Pandya, "A survey on webservice security techniques," in 2018 4th International Conference on Computing Communication and Automation (ICCCA), pp. 1–5, 2018.
- [9] O. B. Fredj, O. Cheikhrouhou, M. Krichen, H. Hamam, and A. Derhab, "An owasp top ten driven survey on web application protection methods," in *Risks and Security of Internet and Systems* (J. Garcia-Alfaro, J. Leneutre, N. Cuppens, and R. Yaich, eds.), (Cham), pp. 235–252, Springer International Publishing, 2021.
- [10] W. Ahmad, Z. Hayat, B. Zafar, F. A. Khan, F. Din, and I. Shah, "A survey on taxonomies of attacks and vulnerabilities in computer systems," *International Journal of Computer Science and Telecommunications*, vol. 3, no. 5, pp. 93–97, 2012.
- [11] S. L. Hansman, "A taxonomy of network and computer attack methodologies," 2003.
- [12] A. C. Panchal, V. M. Khadse, and P. N. Mahalle, "Security issues in iiot: A comprehensive survey of attacks on iiot and its countermeasures," in 2018 IEEE Global Conference on Wireless Computing and Networking (GCWCN), pp. 124–130, 2018.
- [13] C. Simmons, C. Ellis, S. Shiva, D. Dasgupta, and Q. Wu, "Avoidit: A cyber attack taxonomy," in 9th Annual Symposium on Information Assurance (ASIA'14), pp. 2–12, 2014.
- [14] K. F. P. Chan, M. Olivier, and R. P. van Heerden, "A taxonomy of web service attacks," in *Proceedings of the 8th International Conference on Information Warfare and Security: ICIW*, p. 34, 2013.
- [15] M. I. Ladan, "Web services: Security challenges," in 2011 World Congress on Internet Security (WorldCIS-2011), pp. 160–163, IEEE, 2011.
- [16] R. K. K. Meduri, "Webservice security," in *Webservices*, pp. 119–172, Springer, 2019.
- [17] P. Nandan, *Mastering SoapUI*. Packt Publishing Ltd, 2016.
- [18] G. A. Jaafar, S. M. Abdullah, and S. Ismail, "Review of recent detection methods for http ddos attack," *Journal of Computer Networks and Communications*, vol. 2019, 2019.
- [19] M. Jensen, N. Gruschka, and R. Herkenhöner, "A survey of attacks on web services," *Computer Science-Research and Development*, vol. 24, no. 4, pp. 185–197, 2009.
- [20] V. Patel, R. Mohandas, and A. R. Pais, "Attacks on web services and mitigation schemes," in 2010 International Conference on Security and Cryptography (SECRYPT), pp. 1–6, IEEE, 2010.
- [21] A. Singh, A. Sharma, N. Sharma, I. Kaushik, and B. Bhushan, "Taxonomy of attacks on web based applications," in 2019 2nd International Conference on Intelligent Computing, Instrumentation and Control Technologies (ICICT), vol. 1, pp. 1231–1235, 2019.
- [22] A. Boudi, I. Farris, M. Bagaa, and T. Taleb, "Assessing lightweight virtualization for security-as-a-service at the network edge," *IEICE Transactions on Communications*, vol. 102, no. 5, pp. 970–977, 2019.
- [23] A. T. Limited, "Attack tree modeling." <https://www.amenaza.com/>, 2022.
- [24] S. Mehta, G. Raj, and D. Singh, "Penetration testing as a test phase in web service testing a black box pen testing approach," in *Smart Computing and Informatics*, pp. 623–635, Springer, 2018.
- [25] S. Abidi, M. Essafi, C. G. Guegan, M. Fakhri, H. Witt, and H. H. B. Ghezala, "A web service security governance approach based on dedicated micro-services," *Procedia Computer Science*, vol. 159, pp. 372–386, 2019.
- [26] "Web application firewall." https://owasp.org/www-community/Web_Application_Firewall, 2018.
- [27] S. Akshaya, M., and G. Padmavathi, "Taxonomy of security attacks and risk assessment of cloud computing," *Advances in Big Data and Cloud Computing: Proceedings of ICBDC18*. Springer Singapore, 2019.
- [28] S. Yassine, and Y. Maleh, "A systematic review and taxonomy of web applications threats," *Information Security Journal: A Global Perspective* 31.1 (2022): 1–27.
- [29] P. Prinetto, and G. Roascio, "Hardware Security, Vulnerabilities, and Attacks: A Comprehensive Taxonomy." ITASEC. 2020.

- [30] "8 critical web application vulnerabilities and how to prevent them." <https://brightsec.com/blog/web-application-vulnerabilities/>, 2022.



Marcelo Invert Palma Salas holds his Master in Computer Science (2012) and PhD candidate in Computer Science from the University of Campinas (Unicamp), Brazil. Systems Engineer (2005) from the Military School of Engineering. He belongs to the Security and Cryptography Laboratory of the UNICAMP Computer Institute. Professor of the Computer Science Career at the Universidad Mayor de San Andrés, La Paz, Bolivia. Researcher in artificial intelligence (machine learning and deep learning), malware analysis, detection, and classification, and network security over

Tor Network. Cybersecurity specialist, full-stack developer, Linux and Python server administrator. Awards winner of IEEE Richard E. Merwin Award (2013), IEEE Theodore W. Hissey Award (2009), Outstanding Researcher Award from the Bolivian Government (2016). Winner of the FRIDA Grant for the project "Tor Protection against malicious traffic" (2016). IETF 99 Grant in Prague (2017), Cheva Rep. to participate in DOTS communications, Internet Society, etc. IEEE member, TBB (Your Scholarship Bolivia). Reviewer of the journals IEEE Latin America and CLEI, among others.

AUTHORS



Marcelo I. P. Salas

Marcelo Invert Palma Salas holds his Master in Computer Science (2012) and PhD candidate in Computer Science from the University of Campinas (Unicamp), Brazil. Systems Engineer (2005) from the Military School of Engineering. He belongs to the Security and Cryptography Laboratory of the UNICAMP Computer Institute. Professor of the Computer Science Career at the Universidad Mayor de San Andrés, La Paz, Bolivia. Researcher in artificial intelligence (machine learning and deep learning), malware analysis, detection, and classification, and network security over Tor Network. Cybersecurity specialist, full-stack developer, Linux and Python server administrator. Awards winner of IEEE Richard E. Merwin Award (2013), IEEE Theodore W. Hissey Award (2009), Outstanding Researcher Award from the Bolivian Government (2016). Winner of the FRIDA Grant for the project "Tor Protection against malicious traffic" (2016). IETF 99 Grant in Prague (2017), Cheva Rep. to participate in DOTS communications, Internet Society, etc. IEEE member, TBB (Your Scholarship Bolivia). Reviewer of the journals IEEE Latin America and CLEI, among others.

An approach for optimizing resource allocation and usage in cloud computing systems by predicting traffic flow

ARTICLE HISTORY

Received 12 August 2023
Accepted 23 October 2023

Sello Prince Sekwatlakwatla
North-West University
Vanderbijlpark, South Africa
sek.prince@gmail.com
ORCID: 0009-0005-1931-5349

Vusumuzi Malele
North-West University
Vanderbijlpark, South Africa
Vusi.Malele@nwu.ac.za
ORCID: 0000-0001-6803-9030

An Approach for Optimizing Resource Allocation and Usage in Cloud Computing Systems by Predicting Traffic Flow

Sello Prince Sekwatlakwatla
*Department of Computer Science and
 Information Systems
 North-West University
 Vanderbijlpark, South Africa*
 sek.prince@gmail.com
 ORCID: 0009-0005-1931-5349

Vusumuzi Malele
*Department of Computer Science and
 Information Systems
 North-West University
 Vanderbijlpark, South Africa
 Vusi.Malele@nwu.ac.za*
 ORCID: 0000-0001-6803-9030

Abstract—The cloud provides computing resources as a service (scalable and cost-effective storage, management, and accessibility of data and applications) through the Internet. Even though cloud computing offers many opportunities for ICT (information and communication technology), many issues still remain, and the increasing demand for resource management and traffic flow is also becoming increasingly problematic. The amount of data in the cloud computing environment is increasing on a daily basis, which increases data traffic flow. Due to this problem, clients complained about the network speed. Autoregressive Integrated Moving Average (ARIMA), Monte Carlo, Extreme gradient boosting regression (XGBoost), is used in this paper for predicting traffic flow. A Monte Carlo prediction of 84% outperformed ARIMA's prediction of 79.8% and XGBoost's prediction of 71.5%, indicating that Monte Carlo is more accurate than other models when predicting traffic flow in organizational cloud computing systems. A machine learning model will be used for future studies, along with hourly monitoring and resource allocation.

Keywords—Monte Carlo technique, Autoregressive integrated moving average (ARIMA) and Extreme gradient boosting regression

I. INTRODUCTION

As telecommunication networks play an increasingly important role in deploying cloud applications and services in data centers, converged network traffic is growing each year as the cloud computing concept becomes more widely used as a means of providing access to applications and services[1].

Recently, cloud computing has emerged as an innovative way to deliver and host services via the Internet. As a result, business owners are attracted to cloud computing because they do not need to plan ahead for provisioning, and they can start with the fewest resources and expand them only when demand increases[1-3].

Introducing advanced technologies like cloud computing has revolutionized businesses and industries worldwide. It has created a new way of storing and accessing information online[4].Despite offering tremendous opportunities for the IT industry, cloud computing technologies are still in their infancy, with many issues yet to be resolved. It is an increasingly Thoughtful problem for the ICT (Information and Communication Technology) infrastructure organization to

ignore the growing traffic flow and resource management demands.

Managing and maintaining ICT resources is easier and more efficient thanks to cloud computing [6]. The use of cloud computing by many organizations is growing, so if a tool does not exist to forecast cloud computing traffic, resource allocation to clients will be inefficient [8]. Cloud computing is complicated by the inconsistent flow of network traffic, which makes it difficult to predict which network resources will meet the needs of all network clients at a particular moment. Clients complained about slow system times, application timeouts, and high bandwidth usage due to inconsistencies in traffic flow [9].

In recent literature, many cloud computing service providers have been finding it hard to allocate the resources they need to meet their clients' needs, causing system bottlenecks, especially during peak periods [5].

The main contributions of this paper are the simulation experiments conducted and the bibliometric analysis conducted. Monte Carlo, Extreme Gradient Boosting (XGBoost), and Autoregressive Integrated Moving Average (ARIMA) were used to make predictions.

Using a cloud network services traffic flow dataset, this paper compares three prediction techniques used to overcome this challenge. It is structured as follows: an introduction, related work, methods, simulations, results, and conclusions.

II. RELATED WORK

A. Bibliometric analysis of the proposed techniques

For exploring and analyzing large amounts of scientific data, bibliometric analysis has become a popular and rigorous method [1]. As a result, analysis is able to uncover the evolutionary nuances of specific fields. Hence, this paper examines the proposed prediction technique to examine the main key area of this approach.



Fig. 1. Summary of the scopus view

As a result of bibliometric analysis techniques, researchers can gain insights into the selected research area very quickly. Search queries were executed on the Scopus database, and 44 documents were extracted between 2019 and 2023. According to Fig 1, out of 255 authors and 176 author keywords (DE), there is a high prediction rate of 93.43%.



Fig. 2. Most Frequent Words

With reference to the Monte Carlo technique, Autoregressive integrated moving average (ARIMA) and Extreme gradient boosting regression ((XGBoost), Fig 2 presents the word cloud generated using the authors' keywords. Forecasting words are shown in a larger size, such as "machine learning," "learning systems," and "time series" and "time series analysis". Prediction was the area of focus for three proposed models.

B. Gap analysis on the related techniques

A cloud computing service provider must provide adequate and accurate traffic prediction to meet customer needs and support organizations effectively [2-3]. Cloud computing service providers estimate requests and the amount of data they must move frequently to predict system traffic flow. Intelligent transportation cyber-physical cloud control system (ITCPCS) improved accuracy in cloud control traffic; however, it was difficult to predict in real time since the technique is in the development stage [2].

In recent years, Monte Carlo simulation Vibrational mode decomposition (VMD) has been shown to be a more accurate

tool for predicting short-term conditions than long-term models. Limit forecasts to short-term deterministic and probabilistic models [5]. The accuracy of forecasts is tested using publicly available Google trace data and Bit Brain data, but workloads in dynamic clouds may experience substantial forecast errors [13].

There is an overlap between the highest values of subseries, models orders, and weights that are time-dependent and auto-regressive. Auto-regressive integrated moving average models (ARIMAs) [14] produce more accurate forecasts and are more efficient to compute. As most applications and information will have to be moved repeatedly in the future, it is best to move them daily. The activity stream's predictions have been validated in numerous studies [13] and 85% of the predictions have been met. In spite of this, there has been no progress in implementing a framework that predicts traffic using an application framework. Since cloud computing is a continuous process, it is difficult to predict traffic 100% accurately [5-6].

The stationary time series are assumed to remain unchanged over time in this method, which is a Markov chain Monte Carlo approach and a Seasonal Autoregressive Integrated Moving Average (SARIMA) model [6]. In comparison to the existing method, the proposed method has a lower error rate. Artificial Neural Network (ANN) [7] has the ability to predict accurately with a small error or small neurons. It is concise and straightforward to perform Monte Carlo simulations of GP distributions, and errors increase and decrease with sample size [8].

Traffic flow predictions can be enhanced with Multiple Linear Regression Unit-Multi Region Correlation (MRC-MLRU) models [1], which are useful when traffic data is limited. eXtreme Gradient Boosting (XG Boost) [9] this method provides better prediction performance with an R2 of 0.992 and 0.949, respectively. The baseline model tends to be over fitted. This is an improved multivariate linear regression variable parameter spatiotemporal (MLR-VPST) zoning model [11]. Both quality and accuracy have improved, and errors are no longer beyond the process limits [15].

III. METHODS

A. Probabilistic Models

To quantify uncertainty, a probabilistic model integrates the first principle knowledge with data to capture a distribution of state transitions between samples in a batch run by making predictions based on the model predictions. A probability model calculates the probability of certain events occurring rather than monitoring actual data to look for events and data points that conform to a set of rules defined by historical analysis-Monte Carlo.

Monte Carlo technique is a computer-based mathematical technique that accounts for risk quantitatively in forecasting and decision-making [13].

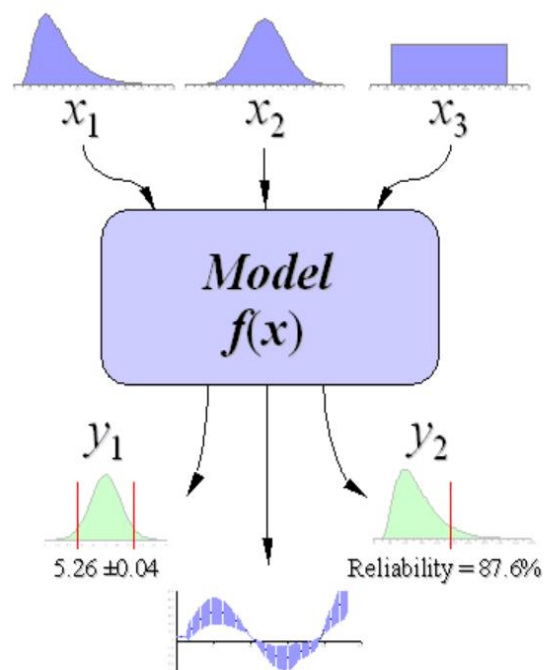


Fig. 3. The simple principle of the Monte Carlo simulation [9]

In order to predict the outcome of a Monte Carlo simulation, the following steps are proposed Fig 3:

Process 1: The first process is to generate a regression parametric.

Technique, $p=f(x_1, x_2, \dots, x_r)$.

Process 2: Set up input generator, $x_{(c1)}, x_{(c2)}, \dots, x_{cq}$

Process 3: Model evaluation and data storage as y_c

Process 4: Process 2 and 3 must be repeated $i=1$ to n

Process 5: Calculate confidence intervals, complete statistics, and histograms based on the results.

B. Machine Learning/Data Mining

In machine learning, data are used to create mathematical models that predict or decide without explicit programming. The algorithms use "training data" to make predictions. Extreme gradient boosting regression (XGBoost), in time series forecasting, models like Extreme Gradient Boosting (XGBoost) can provide reasonable forecasts without tweaking hyper parameters [9]. This is a machine learning algorithm that uses gradient boosting to solve a variety of problems in machine learning.

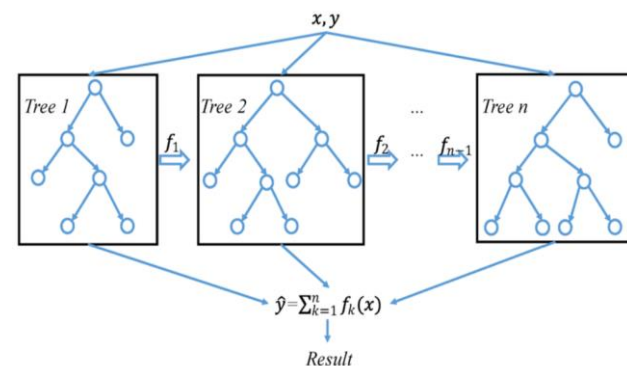


Fig. 4. The simple principle of the XGBoost [11]

It is shown in Fig 4 that XGBoost is an the algorithm moves through iterations, it learns from the residuals of its neighbors. Rather than accepting the majority of the prediction results in Random Forest, it produces a more accurate prediction in this algorithm [1].

$$\hat{y}_l = \sum_{k=1}^n f_k(x_i), \quad f_k \in F, \quad (1)$$

Where f_k In regression trees, the space of trees is defined as,

f_k Assumes the form of a tree, so $f_k(x_i)$ is the outcome of tree k , and $\hat{y}_l$ Evaluate predicted by l request (x_i) ,

C. Statistical Analysis

The Autoregressive integrated moving average (ARIMA) is an advanced analytics technique that uses historical data and statistical models to predict future outcomes. Serial correlation is used to forecast or predict future outcomes with autoregressive integrated moving averages (ARIMA) [4]. In equation 2, the process of calculating autoregressive integrated moving averages (ARIMA) is illustrated by explaining past observations and random errors.

$$K_t = \theta_0 + \varphi_1 k_{t-1} + \varphi_2 y_{t-2} \dots + \varphi_c y_{t-c} + \rho_t + \theta_1 \rho_{t-1} + \theta_2 \rho_{t-2} \dots + \theta_h \rho_{t-h} \quad (2)$$

k_t and ρ_t in Equation. (2) this is a representation of the real value and the error at a particular time interval t , respectively, and $\varphi_1 (i = 1, 2, \dots, p)$ and $\theta_1 (i = 1, 2, \dots, h)$ During auto regression, time series is a linear regression of the p previous values accompanied by the error. Moving average $MA(q)$ describes the current rate of a time series in terms of an error at time t , and q earlier errors.

IV. DATASET

As cloud computing grows, its traffic flow increases every day. As well as uncoordinated traffic signal control and a lack of real-time data, the constant availability of the system is a critical element that cannot be ignored. Nowadays, traffic congestion has a huge impact.

In this study, data were obtained from a South African, Sandton, and automotive industry company with an IT department. As the company outsources other IT infrastructure companies, we are dealing with Cloud computing systems within the organization that are experiencing resource allocation problems. This study

collected comprehensive data from January 1, 2017, to December 31, 2022 (representing daily page loads, unique visitors, new visitors, and returning visitors).

V. SIMULATED SCENARIOS

As illustrated in Fig 5, a three-stage research design is proposed in this study as well as three predictive methods.

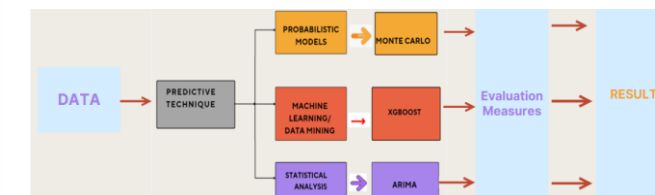


Fig. 5. Proposed technique

The simulation was performed using the Time Series Lab and Crystal Ball.

A model receives data and predicts it, then the model is evaluated and the results are presented in Fig 5.

Method 1 Data > Prediction Techniques > Probabilistic Models > Monte Carlo Techniques > Evaluations > Results.

Method 2 Data > Prediction Techniques > Machine Learning > XGBoost Techniques > Evaluations > Results.

Method 3 Data > Prediction Techniques > Statistical Analysis Techniques > Arima > Evaluations > Results.

VI. RESULTS

A. Evaluation Measures

In evaluation measures, errors between expected and real qualities are defined numerically [15]. The difference between them indicates how well a model did. The data set presented to the network was used to evaluate the performance of the matched network. Mean absolute percentage error (MAPE) was applied to Monte Carlo, ARIMA, and XGBoost techniques in this regard. These measures are defined by equations.

$$MAPE = \frac{1}{N} \sum_{i=1}^N \frac{|S_i - V_i|}{V_i} \times 100 \quad (3)$$

Where S_i Value predicted for observation i ,
 V_i The actual observation value i , and
 N Number of observations.

These prediction analysis techniques have been used to improve cloud network traffic flow predictions. Time Series Lab, Crystal Ball, and real data were used to simulate traffic flow and determine the effectiveness of the proposed model.

B. Monte Carlo technique results

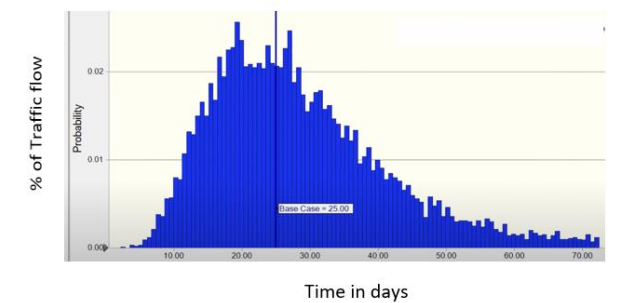


Fig. 6. The normal distribution of the observed data was calculated every 20 seconds according to a 2.2 minute interval.

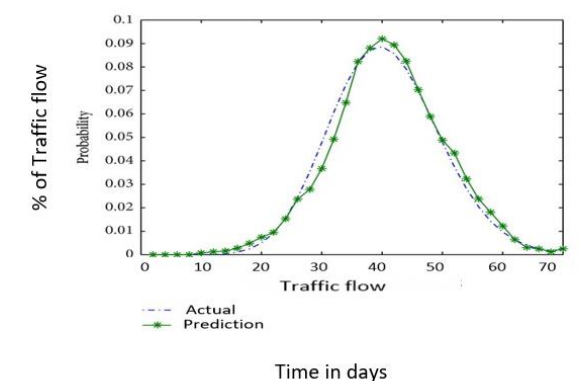


Fig. 7. Observed value

In Fig. 7, when traffic volume was low, the model accuracy increased; however, when traffic volume increased, the model accuracy also decreased, and the prediction was 84%, Fig 7 illustrates Monte Carlo results for the observed data. The normal distribution for the observed data was calculated every 20 seconds according to a 2.2 minute interval as explained in Fig 8. The MAPE prediction errors is 4.49 %.

C. Autoregressive Integrated Moving Averages (ARIMA)

As shown in Fig 8, the results of the Autoregressive Integrated Moving Averages (ARIMA) model were produced within 0.40 seconds, which was 0.40 seconds longer than the previous model. The MAPE prediction error is 5.59 %.

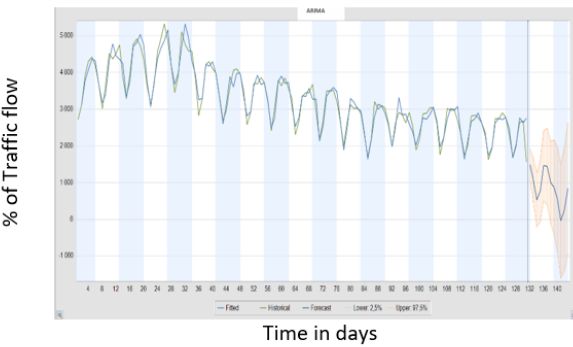


Fig. 8. Autoregressive Integrated Moving Averages (ARIMA)

In Fig. 8, when traffic volume increased, the model accuracy increased; however, when traffic volume decreased, the model accuracy also decreased, and the prediction was 79,8%.

D. Extreme Gradient Boosting Regression (XGBoost)

As shown in Fig 9, extreme gradient boosting (XGBoost) reduced training time to 10 seconds, however, prediction accuracy decreased with a minimum in error evaluation of MAPE 5.88

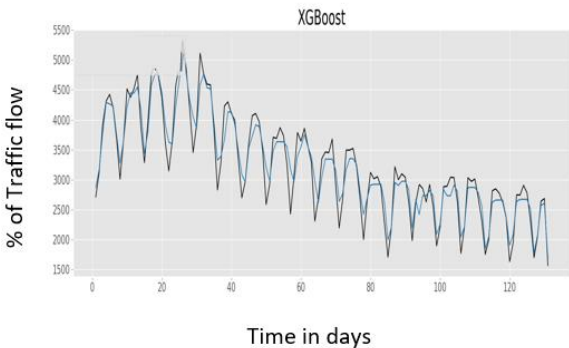


Fig. 9. The XGBoost prediction.

In Fig. 9, when traffic volume increased, the model accuracy increased; however, when traffic volume decreased, the model accuracy also decreased, and the prediction was 71,5%.

VII. DISCUSSION

TABLE I. COMPARISON OF PREDICTION TECHNIQUE ERRORS

Comparison of Techniques Errors			
Evaluation Items	Monte Carlo techniques	ARIMA	XGBoost
MAPE	4.49	4.59	5.88
Training time	~120 obs/sec	~0.40 obs/sec	10.43 obs/sec

In order to identify the best performing techniques, prediction analysis techniques, Monte Carlo techniques, ARIMA techniques, and XGBoost techniques were applied to a real dataset from an online cloud networking application system experiencing traffic congestion.

From 2017 to 2022, a comprehensive set of data was collected on , unique visitors, first-time visitors, and returning visitors on a daily basis.In Table 2, the Monte Carlo training method showed better performance than the ARIMA and XGBoost techniques. It improves prediction accuracy with the minimum amount of errors. Hence, Monte Carlo techniques should be incorporated into traffic prediction models or architectures

VIII.CONCLUSION

The quality of service can only be improved if future workloads are accurately forecasted and resources allocated optimally. By using this predictive analysis, cloud providers can prevent a wide range of losses, such as service outages, excessive or inadequate provisioning of cloud resources, and customer losses.

This paper compares three predictive data analytics techniques to determine how well a cloud computing system predicts future traffic parameters based on their data analytics. A bibliometric analysis review is also conducted to analyses the focus area of the proposed technique, and the findings will contribute to the design of a predictive model for managing cloud computing resources. As a result of traffic control systems' ability to predict future values of traffic parameters, their performance can be improved.

Based on bibliometric analysis, it has been found that the proposed methods are relevant for predicting cloud computing traffic flow, and Monte Carlo techniques are more efficient, than Extreme Gradient Boosting (XGBoost) and Autoregressive integrated moving averages (ARIMA) for analyzing data at a yearly level. Using hourly, daily, weekly, and monthly traffic predictions, future work will determine the most efficient time to allocate cloud computing resources. In addition, other predictive techniques will be included in order to determine which ones will be the most effective in determining the best real-time resource management in cloud computing systems.

A Monte Carlo prediction of 84% outperformed ARIMA's prediction of 79.8% and XGBoost's prediction of 71.5%, indicating that Monte Carlo is more accurate than other models when predicting traffic flow in organizational cloud computing systems. A machine learning model will be used for future studies, along with hourly monitoring and resource allocation.

CONFLICT OF INTEREST

No conflict of interest has been declared by any of the authors.

AUTHOR CONTRIBUTIONS

Ph.D. students developed original drafts after consulting with their supervisors. They conceptualized the methodology, collected the data, prepared the experimental platform, and conducted Time Series Lab and Crystal Ball simulations.

S.P. Sekwatlakwatla, and V. Malele,
“An approach for optimizing resource allocation and usage in cloud computing systems by predicting traffic flow”,
Latin-American Journal of Computing (LAJC), vol. 11, no. 1, 2024.

Following the presentation, supervisors gave significant inputs to the draft, which changed the methodology and provided directions toward developing a traffic analysis model for optimizing resource allocation and utilization in cloud computing systems.

ACKNOWLEDGMENT

The article was written as part of PhD study in Computer Science and Information Systems with Information Technology.

REFERENCES

[1] Zheng, M., Huang, R., Wang, X., Li, X. Do firms adopting cloud computing technology exhibit higher future performance? A textual analysis approach. Journal of International Review of Financial Analysis, 90, (2023) . [Online]. Available: <https://doi.org/10.1016/j.irfa.2023.102866>.

[2] Apat, H. K., Nayak, R., Sahoo, B. A comprehensive review on Internet of Things application placement in Fog computing environment. Journal of Internet of Things, 23,(2023).[Online]. Available:<https://doi.org/10.1016/j.iot.2023.100866>.

[3] Cheng, M., Qu, Y., Jiang, C., Zhao, C. Is cloud computing the digital solution to the future of banking? Journal of Financial stability,63,(2022).[Online]. Available:<https://doi.org/10.1016/j.jfs.2022.101073>.

[4] Luo, J., Gong, Y. Air pollutant prediction based on ARIMA-WOA-LSTM model. journal of Atmospheric Pollution Research,14,(2023). [Online]. Available:<https://doi.org/10.1016/j.apr.2023.101761>.

[5] Jing, J., Magnin, I.E., Frindel, C. Monte Carlo simulation of water diffusion through cardiac tissue models. journal of Medical Engineering and Physics.120,(2023).[Online]. Available: <https://doi.org/10.1016/j.medengphy.2023.104013>

[6] Alés, A., Lanzini, F. Modelling of chemical and magnetic order in Ni-Mn-Al shape memory alloys using Monte Carlo simulations. Journal of Journal of Magnetism and Magnetic Materials, 285,(2023). [Online]. Available:<https://doi.org/10.1016/j.jmmm.2023.171110>

[7] Meerasri , J., Sothornvit, R. Artificial neural networks (ANNs) and multiple linear regression (MLR) for prediction of moisture content for coated pineapple cubes. Journal of Case Studies in Thermal Engineering,33,(2022). [Online]. Available: <https://doi.org/10.1016/j.csite.2022.101942>

[8] Deng, T., Wu, J. Efficient graph neural architecture search using Monte Carlo Tree search and prediction network. Journal of Expert Systems With Applications,213, (2023). [Online]. Available: <https://doi.org/10.1016/j.eswa.2022.118916>

[9] Lee, K., Im, S., Lee, B. Prediction of renewable energy hosting capacity using multiple linear regression in KEPCO system. Journal of Energy Reports,9 , pp.: 343-347 (2023). [Online]. Available:<https://doi.org/10.1016/j.egyr.2023.09.121>.

[10] Afrasiabian, B., Eftekhari, M. Prediction of mode I fracture toughness of rock using linear multiple regression and gene expression programming. Journal of Rock Mechanics and Geotechnical Engineering, 14 , pp.: 1421-1432 (2023). [Online]. Available: <https://doi.org/10.1016/j.jrmge.2022.03.008>.

[11] Silagyi, D.V, Liu, D. Prediction of severity of aviation landing accidents using support vector machine models. Journal of Accident Analysis and Prevention,187, (2023) [Online]. Available:<https://doi.org/10.1016/j.aap.2023.107043>.

[12] Sharma, R., Awasthi, A. An embedded element based 2D finite element model for the strength prediction of mineralized collagen fibril using Monte-Carlo type of simulations. Journal of Biomechanics,108,(2020)[Online]. Available:<https://doi.org/10.1016/j.jbiomech.2020.109867>.

[13] Nadjafi, M., Gholami,P. Probability fatigue life prediction of pin-loaded laminated composites by continuum damage mechanics-based Monte Carlo simulation. Journal of Composites Communications, 32,(2022). [Online]. Available: <https://doi.org/10.1016/j.coco.2022.101161>

[14] He, H., Fan, Y. A novel hybrid ensemble model based on tree-based method and deep learning method for default prediction. journal of Expert Systems With Applications, 176,(2021) .[Online]. Available :<https://doi.org/10.1016/j.eswa.2021.114899>

[15] Santhusitha, D., Karunasingha, K.: Root mean square error or mean absolute error? Use their ratio as well: Information Sciences, 585, pp.: 609-629 (2022). [Online]. Available: <https://doi.org/10.1016/j.ins.2021.11.036>.

AUTHORS



Sello Sekwatlakwatla

A registered PHD student at North West University with over 13 years of IT experience, Sello has expertise in Integration Management, System Implementation, and System Release Management.



Vusumuzi Malele

A senior researcher and Postgraduate supervisor at North West University. An experienced engineer, teacher, research professional and manager with more than 25 years of experience in the ICT industry.

Alzheimer's diagnosis system based on magnetic resonance imaging using the VGG16 algorithm

ARTICLE HISTORY

Received 27 September 2023
Accepted 23 October 2023

Werner Shtaniklao Ucañay Barreto
Universidad Católica Sedes Sapientiae
Lima, Perú
ucanaybarretonsm@gmail.com
ORCID: 0000-0003-3634-7914

Marco Antonio Coral Ygnacio
Universidad Católica Sedes Sapientiae
Lima, Perú
mcoral@ucss.edu.pe
ORCID: 0000-0001-6628-1528

Sistema de Diagnóstico del Alzheimer basado en Imágenes de Resonancia Magnética mediante el Algoritmo VGG16

Alzheimer's Diagnosis System based on Magnetic Resonance Imaging using the VGG16 Algorithm

Werner Shtaniklao Ucañay Barreto
Facultad de Ingeniería
Universidad Católica Sedes Sapientiae
Lima, Perú
ucanaybarretonsm@gmail.com
ORCID: 0000-0003-3634-7914

Marco Antonio Coral Ygnacio
Facultad de Ingeniería
Universidad Católica Sedes Sapientiae
Lima, Perú
mcoral@ucss.edu.pe
ORCID: 0000-0001-6628-1528

Resumen— El diagnóstico temprano del Alzheimer es fundamental para brindar un tratamiento oportuno a los pacientes. En este sentido se ha desarrollado un sistema de diagnóstico del Alzheimer basado en imágenes de resonancia magnética que utiliza un algoritmo de redes neuronales convolucionales denominado VGG16. Se recopiló y procesó imágenes de resonancia magnética de pacientes con y sin Alzheimer. Estas imágenes se utilizaron para entrenar al algoritmo, el cual aprendió a identificar y asociar patrones con la enfermedad. Posteriormente, se realizaron pruebas con un conjunto de imágenes no vistas para evaluar la capacidad de diagnóstico del sistema. Mediante el análisis de las imágenes de resonancia magnética, el algoritmo VGG16 ha demostrado una capacidad superior al 82% para reconocer correctamente dichos signos. Estos resultados validan la efectividad del enfoque basado en inteligencia artificial para el diagnóstico del Alzheimer.

Palabras Clave—diagnóstico del Alzheimer; imágenes de resonancia magnética; VGG16, detección temprana

Abstract— Early diagnosis of Alzheimer's disease is essential to provide timely treatment to patients. In this regard, a system for diagnosing Alzheimer's disease based on magnetic resonance imaging and utilizing a convolutional neural network algorithm called VGG16, has been developed. Magnetic resonance images of patients with and without Alzheimer's disease were collected and processed. These images were used to train the algorithm, which learned to identify and associate patterns with the disease. Subsequently, tests were performed with a set of unseen images to evaluate the diagnostic ability of the system. Through the analysis of magnetic resonance images, the VGG16 algorithm has shown a capacity of over 82% to correctly recognize these signs. These results validate the effectiveness of the artificial intelligence-based approach for diagnosing Alzheimer's disease.

Keywords— Alzheimer's diagnosis; magnetic resonance imaging; VGG16; early detection

I. INTRODUCCIÓN

En los últimos tiempos, el incremento de incidencias de la enfermedad de Alzheimer ha generado una creciente necesidad de detección e intervención temprana [1]. Este trastorno neurodegenerativo afecta a la población mundial y

se ha convertido en una preocupación debido a su impacto negativo en los individuos, sus familias y la sociedad. A menudo se diagnostica cuando ya ha avanzado significativamente, dificultando el manejo y el tratamiento eficaz [2]. Las limitaciones en los enfoques convencionales de diagnóstico, caracterizados por su dependencia de la experiencia y las habilidades de los médicos, pueden dar lugar a variaciones en los resultados y a la posibilidad de que se produzcan errores humanos [3]. Además, los diagnósticos convencionales suelen basarse en análisis subjetivos y pueden verse obstaculizados por el acceso limitado a la información sobre el paciente [4]. Estas limitaciones han impulsado la búsqueda de métodos de diagnóstico más eficaces y precisos. Esta necesidad surge del hecho de que los primeros indicios de enfermedades pueden manifestarse de forma sutil y resultar difíciles de detectar con las técnicas tradicionales [5], lo que provoca retrasos en el diagnóstico y restringe las oportunidades de intervención temprana y tratamiento eficaz.

La posibilidad de un diagnóstico temprano permitirá mejorar la creación de planes de tratamiento eficaces. Los profesionales médicos podrían implementar medidas para frenar la progresión de la enfermedad y mejorar la calidad de vida de los pacientes detectando el Alzheimer en sus primeras fases. La detección temprana también podría reducir la carga económica y emocional de las familias, ya que les da más tiempo para reflexionar y hacer planes para el futuro [6].

Con los avances en técnicas que hacen uso del aprendizaje automático (AM) y la inteligencia artificial (IA), se ha potenciado la capacidad de interpretar y analizar resonancias magnéticas y tomografías por emisión de positrones. Estas tecnologías están siendo aplicadas con el fin de encontrar signos iniciales de la enfermedad de Alzheimer, permitiendo la detección temprana y por ende, la posibilidad de una intervención oportuna [7][8]. El uso de estas tecnologías permite el desarrollo de sistemas de diagnóstico que pueden identificar alteraciones menores en la estructura y funciones cerebrales que pueden predecir el comienzo de la enfermedad

W. S. Ucañay Barreto, and M. A. Coral Ygnacio,
“Alzheimer's Diagnosis System based on Magnetic Resonance Imaging using the VGG16 Algorithm”,
Latin-American Journal of Computing (LAJC), vol. 11, no. 1, 2024.

de Alzheimer, incluso antes de que se manifiesten los indicios clínicos [9]. Sin embargo, la implementación de estas tecnologías avanzadas sigue presentando desafíos significativos, como garantizar su precisión y fiabilidad, así como superar problemas éticos relacionados con la privacidad de los datos y el consentimiento informado [10].

Estos desafíos nos motivan a proponer un sistema de diagnóstico del Alzheimer basado en imágenes de resonancia magnética (IMR) utilizando un algoritmo de redes neuronales convolucionales (VGG16). VGG16 está conformada por 16 capas (capas convolucionales y capas completamente conectadas) en las que aprende a reconocer patrones a partir de una serie de imágenes aplicando filtros y extrayendo características.

El sistema de diagnóstico no se limita únicamente a mejorar la detección temprana de la enfermedad, sino que también aspira a incrementar la calidad de vida de los pacientes al facilitar una intervención más temprana y eficaz. Además, se espera que el sistema ayude a mitigar la carga económica y emocional de las familias al proporcionar diagnósticos más tempranos y precisos.

El sistema se basa en la interpretación de imágenes de resonancia magnética utilizando el algoritmo VGG16. Las imágenes cerebrales fueron obtenidas de The Open Access Series of Imaging Studies (OASIS), un conjunto de datos de imágenes de resonancia magnética de adultos mayores, incluyendo pacientes con y sin Alzheimer. El conjunto de datos contiene imágenes de resonancia magnética estructural y funcional, así como datos clínicos de los pacientes. Estas imágenes son procesadas por el algoritmo, que las compara con patrones previamente entrenados que representan los signos tempranos de la enfermedad. A partir de este análisis, se genera una métrica que cuantifica la similitud con los patrones de la enfermedad. Los resultados han demostrado una capacidad superior al 82% para identificar correctamente los signos iniciales de Alzheimer, validando la efectividad de del enfoque basado en inteligencia artificial.

La estructura del artículo se compone de la siguiente manera: en la sección II, se proporciona un análisis detallado del estado actual del tema, se plantea el problema y se propone una solución precisa para abordarlo. En la sección III, se exponen los resultados derivados de la propuesta, seguidos de un análisis y descripción de los mismos. Finalmente, se presentan las conclusiones extraídas de la investigación y se formulan recomendaciones para futuros trabajos.

II. ESTADO DEL ARTE

Este apartado tiene como objetivo presentar la teoría relevante para entender la propuesta de solución. Aquí, se presenta un panorama de la investigación y desarrollo que se ha hecho en el campo del diagnóstico del Alzheimer basado en imágenes (MRI) y las tecnologías utilizadas en este proceso.

A. Diagnóstico del Alzheimer basado en imágenes (MRI)

El diagnóstico del Alzheimer basado en imágenes (MRI) es un ámbito que ha experimentado un notorio crecimiento en los

últimos tiempos [11]. Esto ha dado lugar al desarrollo de distintos sistemas de diagnóstico del Alzheimer, que gracias a las técnicas y métodos aplicados permiten el reconocimiento de patrones en imágenes cerebrales con el objetivo de detectar signos tempranos de la enfermedad. Su aplicación sirve para la identificación temprana y el seguimiento del Alzheimer, siendo de especial relevancia para el tratamiento y cuidado del paciente.

Los sistemas de diagnóstico del Alzheimer basados en imágenes pueden clasificarse según los recursos utilizados, las características analizadas y la arquitectura del sistema. En cuanto a los recursos utilizados, se pueden emplear tomografía por emisión de positrones (PET), imágenes de resonancia magnética (MRI) o la combinación de ambas técnicas. La MRI es útil para visualizar la estructura cerebral [12], mientras que la PET mide la actividad cerebral y detecta la acumulación de proteínas relacionadas con la enfermedad [13]. En cuanto a las características analizadas, los sistemas pueden centrarse en características morfológicas, como el grosor cortical y el volumen de estructuras cerebrales [14], o en características funcionales, como el flujo sanguíneo cerebral y la actividad metabólica [10]. Algunos sistemas combinan ambas características para mejorar la precisión del diagnóstico. Por último, en función de la arquitectura de la red neuronal, se pueden utilizar redes neuronales convolucionales (especialmente útiles para analizar imágenes) [15], redes neuronales recurrentes (que procesan secuencias de datos y son útiles para analizar imágenes de PET) [16] o redes neuronales adversarias generativas (que generan imágenes de alta calidad y proporcionan información adicional sobre la estructura y función del cerebro) [17].

El enfoque de diagnóstico basado en imágenes de resonancia magnética (MRI) y técnicas de aprendizaje profundo para la enfermedad de Alzheimer ofrece beneficios en múltiples niveles. Desde una perspectiva técnica, ofrece una herramienta de diagnóstico no invasiva y precisa que puede detectar la enfermedad en sus primeras etapas, cuando los cambios estructurales en el cerebro son sutiles y difíciles de identificar a simple vista [18][19]. Esto proporciona a los profesionales de salud una herramienta para enfrentar esta enfermedad compleja.

B. Tecnologías asociadas al diagnóstico del Alzheimer basado en imágenes (MRI)

Los sistemas de diagnóstico del Alzheimer requieren el uso de diversas tecnologías asociadas que permiten el procesamiento y análisis de datos médicos complejos. En consecuencia, resulta esencial examinar la integración de estas tecnologías en los sistemas de diagnóstico del Alzheimer. Un primer punto son las herramientas de procesamiento de imágenes médicas (por ejemplo, Insight Segmentation and Registration Toolkit y Visualization Toolkit) para analizar imágenes y extraer características relevantes [20]. Además, frameworks de aprendizaje profundo como Tensor Flow y PyTorch son esenciales para implementar algoritmos avanzados [10], mientras que bibliotecas de aprendizaje automático como Scikit-learn contribuyen a identificar biomarcadores de la enfermedad [21]. El análisis detallado de

áreas cerebrales específicas se logra mediante software especializado, como FreeSurfer, ANTs y BrainSuite [14]. Por último, se hace mención a las bases de datos de imágenes médicas, ejemplificadas por iniciativas como Open Access Series of Imaging Studies o Alzheimer's Disease Neuroimaging Initiative, que contienen información clínica y permite entrenar y validar el sistema de diagnóstico, lo que mejora su confiabilidad y eficacia en la detección temprana del Alzheimer [22]. Además, es crucial contar con una infraestructura tecnológica adecuada, que incluya computadoras de alta capacidad, unidades de almacenamiento rápidas y tarjetas gráficas de alto rendimiento [23] [24].

C. Métodos de Construcción de sistemas de diagnóstico del Alzheimer

El desarrollo de un sistema de diagnóstico del Alzheimer basado en imágenes (MRI) requiere la comprensión de distintas dimensiones que puede tomar el desarrollo de este sistema, como puede ser la dimensión humana, un punto de vista desde la perspectiva del experto (neurólogo) o en términos del desarrollo del software. Los sistemas de diagnóstico de Alzheimer basados en imágenes son herramientas informáticas especializadas en analizar imágenes médicas para identificar patrones y características asociadas a la enfermedad [25]. Estos sistemas en esencia son software por lo que se construyen mediante diferentes métodos de desarrollo, como los métodos Ágiles [26]. Este enfoque se centra en la flexibilidad, la colaboración y el desarrollo iterativo en función de los comentarios de los usuarios y los requisitos cambiantes. Por otro lado, los sistemas emulan el comportamiento de un neurólogo, quien utiliza evaluaciones visuales, escalas de evaluación y análisis de conectividad para detectar cambios cerebrales relacionados con la enfermedad [27]. Los sistemas de diagnóstico logran emular este comportamiento empleando algoritmos para extraer características relevantes de las imágenes y técnicas de aprendizaje automático para optimizar la precisión del modelo. La implementación puede incluir métodos como el Análisis de componentes principales (PCA), Análisis de textura en regiones de interés, Redes neuronales convolucionales (CNN), Bosques aleatorios (RF), Búsqueda de cuadrícula (Grid Search) y el coeficiente de clustering de Watts-Strogatz [20]. Al relacionar estas dimensiones de manera efectiva, se puede construir un sistema de diagnóstico del Alzheimer que mejore la precisión y la robustez de las predicciones utilizando información y modelos complementarios.

D. Problemas asociados a la implementación de Sistemas de diagnóstico del Alzheimer

La implementación de un sistema de diagnóstico del Alzheimer basado en imágenes (MRI) también tiene sus desafíos. Estos problemas pueden ser técnicos, metodológicos o éticos.

Se presentan desafíos significativos relacionados con problemas técnicos en la implementación de estas tecnologías avanzadas. Estos problemas incluyen la calidad de las imágenes, ya que estas deben ser lo suficientemente detalladas para permitir un análisis preciso, y deben ser normalizadas para asegurar que las comparaciones entre las imágenes sean válidas [28]. Además, el manejo de grandes volúmenes de datos puede ser un desafío, ya que el análisis de imágenes de MRI requiere un gran almacenamiento y poder de procesamiento [29]. Los problemas metodológicos suelen

estar relacionados con la selección y aplicación de los métodos de análisis. La elección de los parámetros adecuados para los algoritmos de aprendizaje automático puede ser un proceso complicado, y los resultados pueden variar significativamente dependiendo de estos parámetros [28]. Además, la validación de los resultados es un paso crucial, y el uso de conjuntos de datos de prueba independientes es esencial para asegurar que los modelos son generalizables y no simplemente sobre ajustados a los datos de entrenamiento [22]. Finalmente, los problemas éticos son una consideración importante en el desarrollo de cualquier sistema médico. En el caso del diagnóstico del Alzheimer basado en imágenes (MRI), estos problemas pueden incluir cuestiones de privacidad y consentimiento, ya que las imágenes cerebrales son datos sensibles y personales. Además, el diagnóstico temprano de la enfermedad de Alzheimer puede tener implicaciones psicológicas significativas para los pacientes, y debe manejarse de manera cuidadosa y ética [3].

III. EL PROBLEMA DEL DIAGNÓSTICO DEL ALZHEIMER BASADO EN IMÁGENES

El diagnóstico de la enfermedad de Alzheimer plantea importantes retos a la comunidad médica, sobre todo cuando se basa en técnicas de imagen como la resonancia magnética (MRI). Estos problemas se derivan de la variabilidad de las imágenes cerebrales, la presencia de rasgos sutiles o inespecíficos del Alzheimer, la ausencia de indicadores definidos [30] y las desafíos que experimentan los expertos a la hora de interpretar las imágenes [31].

En primer lugar, la variabilidad de las imágenes cerebrales plantea un reto importante en el diagnóstico de la enfermedad de Alzheimer. Los cerebros humanos varían considerablemente de una persona a otra, tanto en términos de estructura física como de función cognitiva [32]. Esta variabilidad natural puede dificultar la distinción entre los cambios normales relacionados con la edad y los cambios patológicos asociados a la enfermedad de Alzheimer. Además, diversos factores como la edad, el sexo, la educación y la genética pueden influir en la estructura y la función cerebrales [33], complicando aún más la tarea del diagnóstico. Por ejemplo, algunas investigaciones han demostrado que el cerebro de las mujeres tiende a envejecer de forma diferente al de los hombres, lo que puede influir en las características estructurales visibles en una resonancia magnética. Del mismo modo, se ha demostrado que el nivel educativo de una persona influye en la estructura cerebral, sobre todo en las áreas relacionadas con el lenguaje y la cognición [34]. A nivel genético, algunos genes pueden afectar a la estructura y función del cerebro, lo que aumenta aún más la variabilidad observada en las imágenes cerebrales [16]. Esta amplia variabilidad dificulta una definición precisa entre lo que se considera normal y lo que indica la enfermedad de Alzheimer.

Además de la variabilidad de las imágenes cerebrales, la detección de características leves de la enfermedad de Alzheimer en las imágenes dificulta aún más el procedimiento de diagnóstico. La enfermedad de Alzheimer se distingue por la degeneración progresiva del tejido cerebral, sobre todo en el hipocampo y otras zonas del cerebro relacionadas con la memoria [35]. Sin embargo, estos cambios suelen ser sutiles y pueden no ser perceptibles hasta que la enfermedad ha progresado significativamente. Además, esta atrofia no es exclusiva de la enfermedad de Alzheimer; otras formas de demencia pueden presentar patrones similares de atrofia

cerebral, por lo que resulta difícil diferenciarlas basándose únicamente en los resultados de las pruebas de imagen. Para añadir mayor complejidad, existen características inespecíficas como las hiperintensidades de la sustancia blanca [36]. Se trata de áreas de mayor intensidad de señal que se observan en las resonancias magnéticas y que pueden encontrarse tanto en cerebros que envejecen normalmente como en los que padecen la enfermedad de Alzheimer. Estas características, aunque indican un posible daño vascular, no son específicas del Alzheimer, lo que aumenta aún más la complejidad del diagnóstico.

Un desafío significativo en el diagnóstico por imagen del Alzheimer radica en la carencia de marcadores claros y definitivos. Aunque la presencia de placas de beta-amiloide y ovillos neurofibrilares son señales distintivas de la enfermedad de Alzheimer [37], estas alteraciones patológicas no son directamente observables mediante las técnicas convencionales de resonancia magnética. Métodos de imagen avanzados como la tomografía por emisión de positrones (PET) pueden visualizar las placas amiloides, pero estas técnicas están menos disponibles y exponen a los pacientes a radiaciones ionizantes. Por lo tanto, la resonancia magnética, que no es invasiva, sigue siendo la principal modalidad de imagen para evaluar a los pacientes con sospecha de enfermedad de Alzheimer, a pesar de sus limitaciones.

Por último, la interpretación de las imágenes cerebrales por parte de los expertos médicos representa un desafío significativo. Se requiere un alto grado de experiencia para identificar e interpretar con precisión los cambios sutiles asociados a la enfermedad de Alzheimer. Los sutiles cambios que se producen en el cerebro debido a la enfermedad de Alzheimer no son fáciles de identificar y pueden pasar desapercibidos o malinterpretarse con facilidad [33], lo que conduce a un diagnóstico erróneo. Además, la interpretación puede ser muy subjetiva, y distintos profesionales sanitarios pueden interpretar la misma imagen de formas diferentes [7][38], un fenómeno conocido como variabilidad entre evaluadores. Esta incoherencia puede llevar a confusión y a un diagnóstico potencialmente incorrecto.

A pesar de estos retos, el diagnóstico por imagen de la enfermedad de Alzheimer presenta ciertas ventajas sobre los métodos de diagnóstico tradicionales. Los métodos tradicionales suelen basarse en pruebas cognitivas, que pueden verse influidas por diversos factores, como el estado de ánimo del paciente, su educación y sus antecedentes culturales [19]. En cambio, el diagnóstico por imagen proporciona una medida más objetiva de la estructura y el funcionamiento del cerebro. Las imágenes también pueden detectar cambios en el cerebro que se producen antes de la aparición de los síntomas cognitivos [12], lo que potencialmente permite un diagnóstico y una intervención más tempranos. Además, los avances tecnológicos y los algoritmos de aprendizaje automático pueden mejorar la precisión y la eficacia de la interpretación de imágenes [39], mitigando así algunos de los problemas asociados a la interpretación humana.

IV. PROPUESTA DE DIAGNÓSTICO DEL ALZHEIMER BASADO EN IMÁGENES

El Instituto Nacional de Ciencias Neurológicas, especializado en Neurología y Neurocirugía [40], se enfrenta a importantes dificultades en la detección temprana del Alzheimer debido a los exhaustivos requisitos de las evaluaciones de laboratorio y de imagen cerebral. El diagnóstico del deterioro cognitivo leve, que requiere pruebas especializadas de las que actualmente no se dispone en el país, complica aún más el proceso, lo que a menudo conlleva un retraso en la detección.

En respuesta a esta problemática, se propone un sistema de diagnóstico de la enfermedad de Alzheimer que utiliza imágenes de resonancia magnética (MRI) junto con un algoritmo de red neuronal convolucional (CNN). Los datos de entrada incluyen resonancias magnéticas, obtenidas de la base de datos de Oasis-1 [41], que proporcionan información sobre la estructura del cerebro. OASIS-1 es un conjunto de datos de imágenes cerebrales de 416 pacientes de entre 18 a 96 años, incluye imágenes de resonancia magnética ponderadas en T1 además de datos clínicos y cognitivos. El sistema se basará en los métodos utilizados por los neurólogos para identificar un posible desarrollo de la enfermedad, como puede ser una búsqueda de alteraciones específicas, como el agrandamiento de los ventrículos cerebrales, la pérdida de volumen del hipocampo, los cambios en la sustancia blanca y el adelgazamiento de la sustancia gris, lo cual puede ser indicativo de la progresión de la enfermedad de Alzheimer.

Es importante destacar que este sistema se propone como una herramienta complementaria en el diagnóstico médico. Trabajaría en conjunto con los profesionales de la salud, respaldando su experiencia y conocimientos clínicos en la detección y diagnóstico de la enfermedad de Alzheimer. La integración de esta herramienta tiene como objetivo reducir los retrasos en la detección y mejorar la precisión diagnóstica en beneficio de los pacientes.

A. Marco Teórico

En esta sección, se explorarán la teoría y conceptos clave relacionados con la arquitectura de la una red convolucional, siendo más específico, el algoritmo VGG16.

1) Capas convolucionales

Una capa convolucional es un componente fundamental en las redes neuronales convolucionales. Funciona mediante la aplicación de filtros, llamados 'kernels', que recorren toda la imagen y generan mapas de características, resaltando los detalles importantes para el análisis como se muestra en la Figura 1. Estos filtros permiten detectar patrones locales en las imágenes, tales como bordes, texturas y colores, lo que facilita la identificación de características útiles para tareas como la clasificación y detección de objetos. Este filtro es una pequeña matriz de pesos. Al mover el filtro por la imagen (de izquierda a derecha, de arriba abajo), la red es capaz de aprender patrones locales, como bordes o texturas. En el modelo VGG16, la arquitectura incluye múltiples capas convolucionales con filtros pequeños (3x3), esto permite al modelo aprender patrones más complejos.

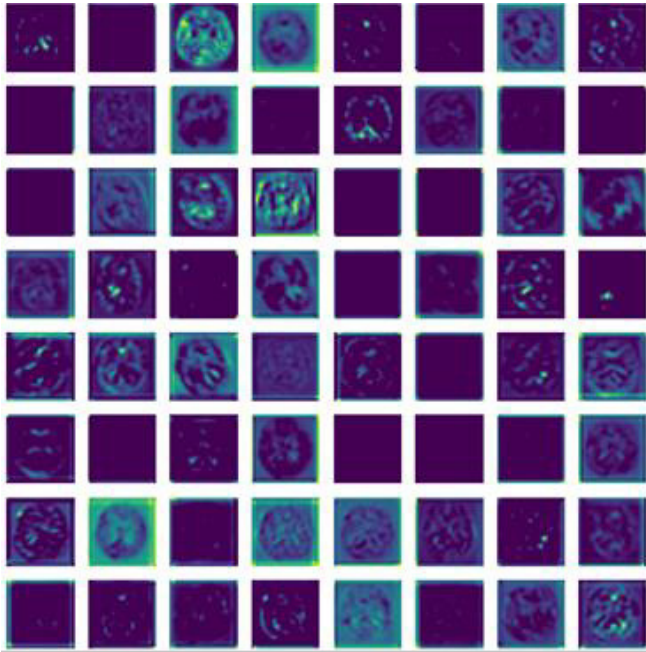


Fig. 1. Filtros convolucionales aplicados a una imagen para resaltar los bordes horizontales

2) Capas de Pooling

Después de las capas convolucionales, se suelen utilizar capas de agrupamiento para reducir el tamaño espacial (anchura y altura) de la característica convolucionada. Esto sirve para disminuir la potencia de cálculo necesaria para procesar los datos mediante la reducción de la dimensionalidad. Existen varios tipos de operaciones de pooling, pero el más común es el Max pooling. La agrupación máxima selecciona el valor máximo de una región del mapa de características cubierta por el filtro. La capa de pooling en VGG16 utiliza 2x2 Max pooling, lo que significa que mira en una cuadrícula de 2x2 del mapa de características y elige el valor más grande, reduciendo efectivamente el tamaño del mapa de características a la mitad.

La Figura 2 muestra una sección de la imagen de la capa de pooling. La capa de pooling reduce el tamaño de la imagen de entrada, y mantiene la información más relevante. En este caso, la capa de pooling utiliza un tamaño de ventana de 2x2 y un stride de 2. Esto significa que cada región de 2x2 píxeles en la imagen de entrada se reduce a un solo píxel en la imagen de pooling.

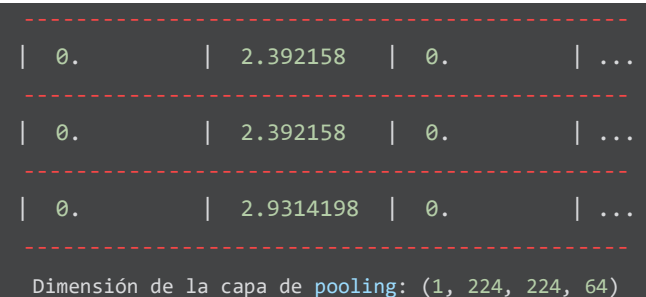


Fig. 2. Resultado de la capa pooling aplicado a 3 pixeles

La forma de la matriz permite comprender la configuración espacial de los resultados de la agrupación: la primera dimensión indica el tamaño del lote (uno, para una sola imagen procesada), y las dimensiones siguientes representan la altura, la anchura y la profundidad (número de

canales) del resultado de la agrupación. Estas dimensiones ayudan a comprender la estructura y la representación de las características extraídas. Cada valor en la imagen de pooling representa el valor máximo de la región de 2x2 píxeles correspondiente en la imagen de entrada. Por ejemplo, el valor 2.392158 en la fila 1, columna 1 de la imagen de pooling representa el valor máximo de la región de 2x2 píxeles correspondiente en la imagen de entrada.

Los valores cercanos a cero en la matriz sugieren que se seleccionaron características insignificantes durante la agrupación, ya sea debido a la selección de valores mínimos o a la falta de características significativas en esa zona. Por otro lado, los valores distintos de cero expresan la existencia de rasgos notables en esa región, lo que representa datos cruciales recuperados y condensados por la capa de agrupación.

B. Desarrollo del Sistema de diagnóstico del Alzheimer basado en imágenes

Para el sistema propuesto, se empleará el modelo VGG16, un modelo de aprendizaje profundo preentrenado, usado para tareas de clasificación de imágenes gracias a su capacidad para identificar patrones en conjuntos de datos complejos, como los derivados de imágenes de resonancia magnética. Al adaptar el modelo VGG16 a nuestros requisitos específicos, el sistema está diseñado para diagnosticar la enfermedad de Alzheimer mediante el examen de patrones en las imágenes de resonancia magnética de los pacientes.

Este modelo (Figura 3) se caracteriza por su profundidad y relativa sencillez. Compuesto por 16 capas, incluye capas convolucionales, capas de agrupación y capas totalmente conectadas. Cada capa convolucional está repleta de múltiples filtros diseñados para extraer diversas características de las imágenes a distintos niveles de abstracción. Las capas de agrupación sirven para reducir la dimensionalidad de las características obtenidas, reteniendo sólo las más significativas.

A continuación, las capas totalmente conectadas se encargan de la clasificación final de las imágenes, segregándolas en distintas categorías. Esta arquitectura resalta la capacidad del modelo VGG16 para proporcionar interpretaciones detalladas y matizadas de las imágenes, lo que lo hace adecuado para los objetivos de nuestro sistema.

El primer paso consiste en entrenar el modelo utilizando la base de datos Oasis-1 [41], la cual es una base de datos ampliamente utilizada en la investigación del Alzheimer. Oasis-1 contiene imágenes de resonancia magnética (RM) de cerebros de individuos sanos y con diferentes etapas de la enfermedad de Alzheimer.

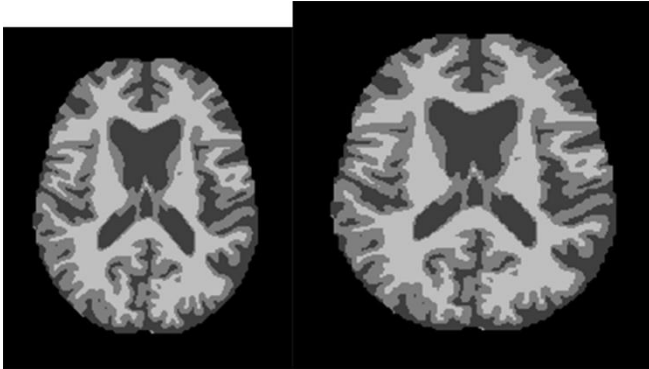


Image		Image	
Dimensions	176 x 208	Dimensions	224 x 224
Width	176 pixels	Width	224 pixels
Height	208 pixels	Height	224 pixels
Bit depth	8	Bit depth	32
File		File	
Name	OAS1_0307_MR1_mpr_n4_anon_111...	Name	OAS1_0307_MR1_mpr_n4_anon_111...
Item type	GIF File	Item type	PNG File

Fig. 4. Imagen original (izq.) y resultado de la normalización (der.)

Las imágenes de resonancia magnética (RM) de esta base de datos se preprocesan para adaptarlas a los requisitos de la arquitectura VGG16. Como se muestra en la Figura 4, estos requisitos implican redimensionar las imágenes a dimensiones de 224x224 píxeles. En la Ecuación (1), se aplica la normalización de las imágenes, la cual utilizó la normalización del valor del píxel, un tipo de normalización lineal:

$$PxN = \frac{PV - PMin}{PMax - PMin} \quad (1)$$

Donde:

PxN = Pixel Normalizado
 PV = Valor del Pixel
 $PMax$ = Valor Máximo del Pixel
 $PMin$ = Valor Mínimo del Pixel

Este algoritmo normaliza los valores de los píxeles en el rango de 0 a 1, lo que hace que las imágenes sean más comparables entre sí. Luego se realiza la preparación de los datos en el formato adecuado para el entrenamiento, como es el formato de extensión PNG.

Finalmente, las imágenes son divididas en cuatro categorías (sin Alzheimer, leve, moderado y severo) para su entrenamiento y validación. De la base de datos Oasis-1 Se seleccionaron 400 pacientes, que se dividieron en un conjunto de entrenamiento de 296 pacientes y un conjunto de validación de 104 pacientes, lo que resulta en un modelo preentrenado capaz de clasificar las imágenes de acuerdo con las clases establecidas y los pesos asignados.

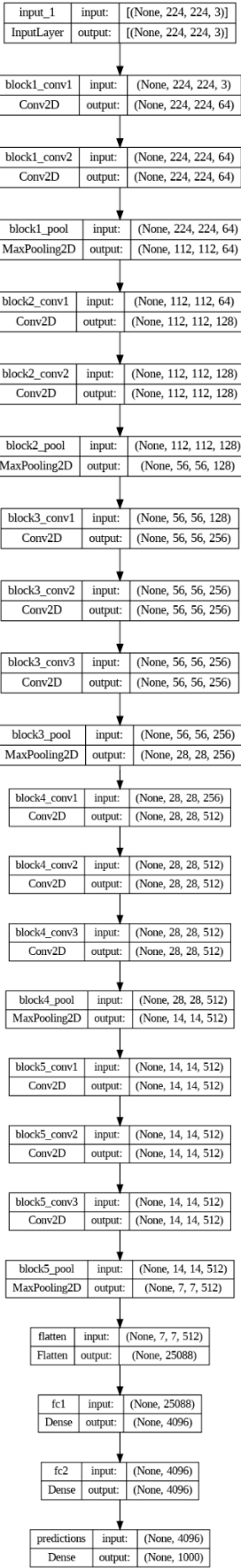


Fig.3. Modelo de VGG16

C. Aplicación

Esta subsección se adentra en el proceso de entrenamiento y procesamiento de las imágenes. Además, se detalla el resultado del entrenamiento y el proceso que sigue las imágenes a través de la red.

1) Entrenamiento

En el entrenamiento, la red neuronal aprende a reconocer patrones y características específicas en las imágenes cerebrales que pueden indicar la presencia de la enfermedad de Alzheimer. Se utilizaron un total de 1,305 imágenes para el conjunto de entrenamiento y 458 imágenes para el conjunto de validación. Estas imágenes representan una variedad de casos de pacientes con y sin Alzheimer, lo que proporciona una muestra diversa y representativa para entrenar y evaluar el modelo.

El conjunto de entrenamiento de 1,305 imágenes se utilizó para ajustar los pesos y los parámetros de la red neuronal, permitiendo que el modelo aprenda a reconocer las características distintivas asociadas con el Alzheimer en las imágenes de MRI. Por otro lado, el conjunto de validación de 458 imágenes se utilizó para evaluar el rendimiento del modelo durante el entrenamiento y evitar el sobreajuste. Esta separación de los conjuntos de entrenamiento y validación ayuda a medir la capacidad del modelo para generalizar y realizar predicciones precisas en datos no vistos previamente.

Durante el proceso de entrenamiento, se obtuvo una precisión del modelo superior al 82%, lo que indica una gran capacidad de clasificación y reconocimiento de las características distintivas asociadas con el Alzheimer en las imágenes de MRI. Esta destacada precisión se evidencia en la Figura 5 "Precisión del Modelo", donde se puede apreciar el progreso del modelo a medida que se entrenaba.

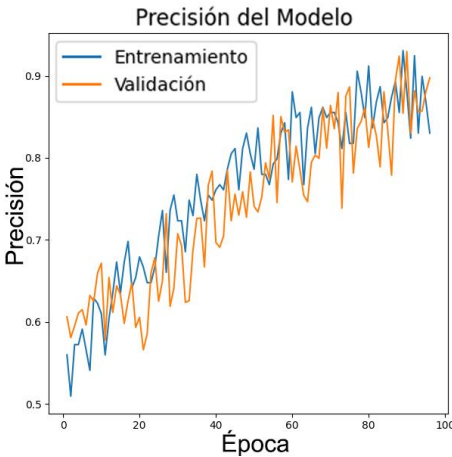


Fig. 5. Gráfica de Precisión del Modelo

Por otro lado, la pérdida del modelo, que se refiere a la discrepancia entre las predicciones del modelo y los valores reales durante el entrenamiento, se ha mantenido en un nivel muy bajo, alrededor del 0.5%. Esta baja pérdida indica que el modelo ha aprendido de manera efectiva a generalizar y ajustarse a los datos de entrenamiento, lo que a su vez contribuye a su alta precisión. Figura 6 "Pérdida del Modelo" muestra una tendencia decreciente a medida que avanza el entrenamiento, lo que demuestra la mejora constante y la capacidad del modelo para reducir los errores en sus predicciones.

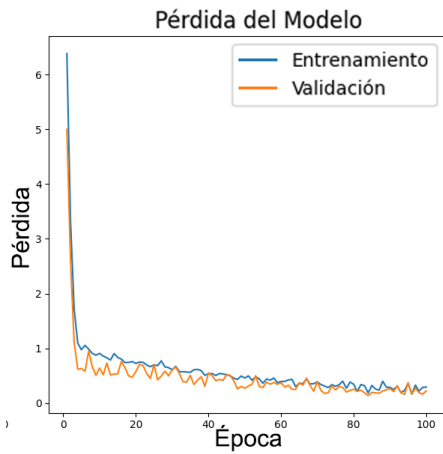


Fig. 6. Gráfica de Pérdida del Modelo

Las Figuras 5 y 6 ayudan a evaluar el rendimiento y el progreso del entrenamiento de la red VGG16 en el sistema de diagnóstico del Alzheimer, ya que permiten visualizar claramente cómo evolucionaron la precisión y la pérdida del modelo a lo largo del entrenamiento y cómo se alcanzaron los niveles de precisión y la baja pérdida mencionados anteriormente. Por otro lado, la figura 6 presenta una representación visual de los patrones observados en imágenes de resonancia magnética en diferentes categorías de Alzheimer: Sin Alzheimer, leve, medio y moderado. Se ha utilizado un enfoque basado en el análisis de agrandamiento de los cuencos cerebrales y adelgazamiento de la materia blanca para identificar las características distintivas asociadas con cada categoría.

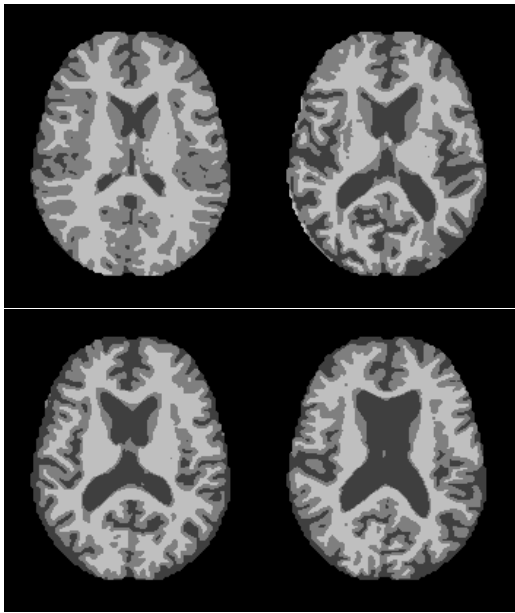


Fig. 7. Patrón de agrandamiento de los cuencos cerebrales y adelgazamiento de la materia blanca en diferentes categorías de Alzheimer

La Figura 7., obtenida de la base de datos Oasis, ilustra cómo estos patrones cambian a medida que el Alzheimer progresa, lo que proporciona una comprensión más clara de las alteraciones estructurales en el cerebro relacionadas con esta enfermedad neurodegenerativa.

Una vez entrenada la red, se realiza el procesamiento de las imágenes para obtener su diagnóstico basándose en los patrones antes mencionados. Para ello, en primer lugar, se

importan las bibliotecas necesarias para ejecutar el programa. Estas librerías incluyen numpy para operaciones numéricas, nibabel para leer imágenes NIfTI, tensorflow para construir y entrenar el modelo de red neuronal, OpenCV para procesamiento de imágenes, e imutils para funciones adicionales de procesamiento de imágenes.

Luego se definen las constantes como la ruta a los pesos del modelo preentrenado y las dimensiones de las imágenes, así como la ruta de la imagen para realizar la predicción del modelo.

A continuación, se realiza la carga del diccionario de clases y preprocesamiento de la imagen para el modelo. El diccionario de clases asigna los índices de las clases a sus nombres. A continuación, se carga la imagen previamente guardada, la normaliza y añade una dimensión adicional para que coincida con la forma de entrada del modelo.

El modelo se construye utilizando la arquitectura VGG16, elimina las capas superiores y añade capas personalizadas. Se cargan los pesos del modelo preentrenado y se realiza la predicción.

La predicción se postprocesa seleccionando la clase con mayor probabilidad. A continuación, el nombre de la clase predicha y su correspondiente probabilidad se escriben en la imagen original, que se guarda en un directorio correspondiente a la clase predicha.

```
Console output:
Predicting
C:\xampp\htdocs\Prototipo\src\static\imagenes\OAS1_0001_MR1_mpr_n4_anon_111_t88_masked_gfc_fseg.nii
{'mild_dementia': 0, 'moderate_dementia': 1, 'no_dementia': 2, 'slight_dementia': 3}
1/1 [=====] - ETA: 0s
1/1 [=====] - 0s 40ms/step
[0]

mild_dementia - [[0.8038744 0.00100118 0.30867776 0.13876037]]
```

Fig. 8. Salida después de procesar la imagen

En la Figura 8, las cifras entre corchetes simbolizan las probabilidades vinculadas a las predicciones del modelo, con cuatro categorías en este caso: Sin Alzheimer, Alzheimer leve, moderado y moderado/intermedio. Cada número denota la probabilidad de que la imagen analizada pertenezca a cada categoría. Por ejemplo, en la predicción "demencia leve - [[0.803 0.001 0.308 0.138]]", los valores [0.803 0.001 0.308 0.138] sugieren que la imagen tiene alrededor de un 80% de probabilidades de pertenecer a la categoría "Alzheimer moderado", un 0,1% de probabilidades de "Alzheimer moderado/intermedio", un 30% de probabilidades de "No Alzheimer" y un 13% de probabilidades de "Alzheimer leve".

2) Mapas de Calor

Con el objetivo de identificar características importantes para el diagnóstico del Alzheimer, se obtienen mapas de calor para las imágenes de resonancia magnética del cerebro (MRI). Se comienza con imágenes preprocesadas y ya cargadas en el

modelo. El proceso se enfoca en la generación de mapas de calor que destacan las áreas de interés en la imagen. Una vez que la imagen preprocesada se encuentra en el modelo VGG16, el objetivo es extraer información relevante de la última capa convolucional. La Figura 9, muestra el resultado de la última capa convolucional (altamente especializada), que contiene características de alto nivel que el modelo ha aprendido a reconocer.

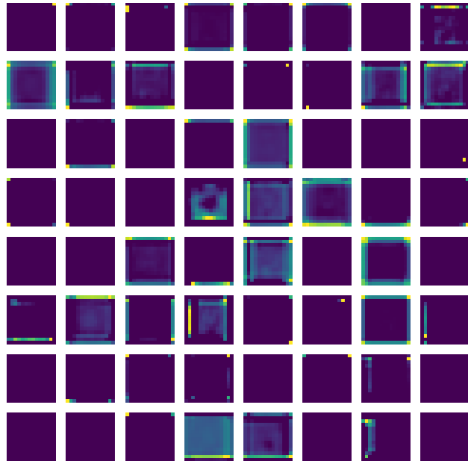


Fig. 9. Última capa convolucional del algoritmo VGG16. La imagen muestra un mapa de características de 512 canales, cada uno de los cuales representa una característica diferente en la imagen.

Una vez obtenida la capa, se multiplica la salida de esta capa por sus respectivos pesos. Los pesos representan la importancia de cada característica aprendida por la red. Luego, se suman estos resultados a lo largo de la capa para obtener un mapa de características ponderado.

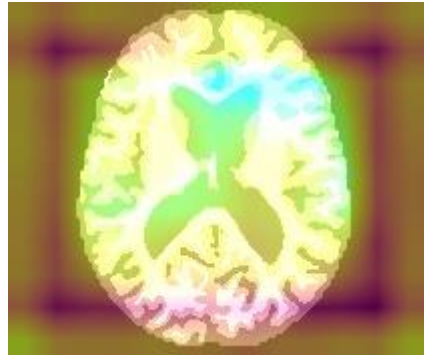


Fig. 10. Mapa de calor obtenido

En la Figura 10, se observa el mapa de calor resultante. Se ha asignado un tono celeste a las áreas en las que se ha centrado más la atención del modelo. Esta coloración indica la ubicación de posibles cambios en la materia blanca del cerebro, que están relacionados con la enfermedad de Alzheimer. Las áreas resaltadas en rojo y amarillo representan niveles crecientes de intensidad o importancia en relación con los cambios estructurales detectados. Estos colores resaltan las regiones que podrían ser de mayor relevancia para la detección y comprensión de las alteraciones en la estructura cerebral asociadas con el Alzheimer.

V. RESULTADOS

Se aplicó el modelo a dos casos: (i) un paciente sano; y (ii) otro paciente con Alzheimer en estado leve. En cada experimento, se utilizaron todos los patrones aprendidos por la red neuronal. Sin embargo, en el primer experimento, se le dio un mayor peso al patrón de cambios en la materia blanca, mientras que, en el segundo experimento, se le dio un mayor peso al patrón de cambios en los cuencos cerebrales. Se utilizó un conjunto de datos de imágenes no vistas previamente para la experimentación.

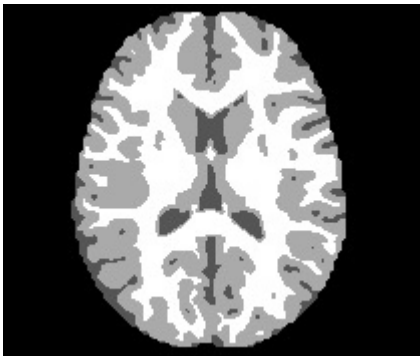


Fig. 11. Resonancia magnética cerebral de un paciente sano

```
Console output: Analisis de imagen MRI
Fecha y Hora: 2023-06-14 02:01 PM
Imagen Procesada:
C:\Users\Ygnacio\Desktop\imagenes_prueba\FSL_SEG\OAS1_0006_MRI_mpr_n4_anon_111_t88_masked_gfc_fseg.nii
Carga de imagen MRI completa.
Procesando imagen...
Extracción de características utilizando el modelo VGG16...
1/1 [=====] - ETA: 0s
1/1 [=====] - 0s 287ms/step
Progreso: 50% completado
1/1 [=====] - ETA: 0s
1/1 [=====] - 0s 51ms/step
Progreso: 100% completado
Resultados: Sin Alzheimer
Biomarcadores Asociados:
No se hallaron biomarcadores
```

Fig. 12. Resultado del procesamiento y prediccion del algoritmo en un paciente sano

Para el experimento (i) se utilizó la Figura 11 como imagen de entrada, la cual reveló resultados negativos para la presencia de Alzheimer (Figura 12). En este caso, no se encontraron biomarcadores característicos de la enfermedad en dicha imagen. Para visualizar y resaltar posibles cambios en la materia blanca del cerebro, se utilizó un mapa de calor. En este mapa, se enfocó especialmente en las zonas donde se podrían encontrar cambios estructurales relevantes. La imagen de prueba de entrada utilizada en el experimento proporciona una representación visual del estado de la muestra analizada. Esta imagen, obtenida a partir de resonancia magnética (MRI), fue seleccionada para evaluar la presencia o ausencia de características específicas asociadas con la condición bajo estudio.

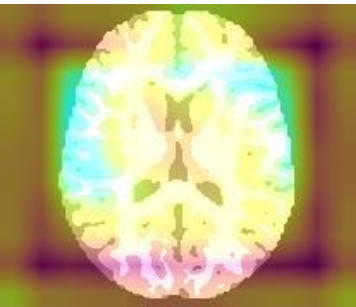


Fig. 13. Mapa de Calor resultante del experimento (i)

La figura 13 muestra el mapa de calor de la imagen procesada. Se asignó un color celeste a las áreas en las que se realizó un mayor enfoque, lo que indica la ubicación de posibles cambios de la materia blanca asociados con el Alzheimer. Por otro lado, las áreas resaltadas en rojo y amarillo representan niveles crecientes de intensidad o importancia en relación con los cambios estructurales detectados. Estos colores resaltan las regiones que podrían ser de mayor relevancia para la detección y comprensión de las alteraciones en la estructura cerebral asociadas con el Alzheimer.

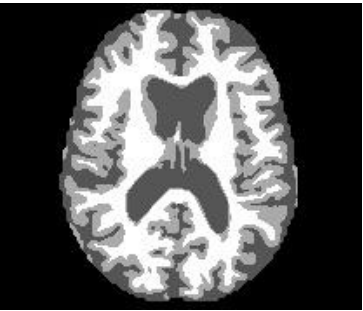


Fig. 14. Resonancia magnética cerebral de un paciente con Alzheimer en estado Leve

```
Imagen Procesada:
C:\Users\Ygnacio\Desktop\imagenes_prueba\FSL_SEG\OAS1_0010_MRI_mpr_n4_anon_111_t88_masked_gfc_fseg.nii
Carga de imagen MRI completa.
Procesando imagen...
Extracción de características utilizando el modelo VGG16...
1/1 [=====] - ETA: 0s
1/1 [=====] - 0s 194ms/step
Progreso: 50% completado
1/1 [=====] - ETA: 0s
1/1 [=====] - 0s 42ms/step
Progreso: 100% completado
Resultados: Alzheimer Leve
Biomarcadores Asociados:
-Aumento de la presencia de lesiones cerebrales focales.
-Disminución del tamaño del cerebro en general.
-Aumento del grosor de la corteza cerebral en regiones específicas.
```

Fig. 15. Resultado del procesamiento y prediccion del algoritmo en un paciente con Alzheimer en estado Leve

Para el experimento (ii) se utilizó la Figura 14 como imagen de entrada la cual arrojó resultados positivos para la presencia de Alzheimer en una etapa leve (Figura 15). El algoritmo empleado se enfocó especialmente en el análisis de los cuencos cerebrales, ya que se ha demostrado que estas estructuras pueden verse afectadas en etapas tempranas de la enfermedad. Se utilizó una imagen representativa del objeto de estudio para llevar a cabo el análisis y las pruebas. La imagen seleccionada fue sometida a diferentes técnicas de procesamiento y análisis para extraer información significativa y revelar patrones o propiedades relevantes. El uso de esta imagen en el experimento proporcionó una base para realizar mediciones, comparaciones y conclusiones objetivas, contribuyendo así al avance de la investigación.

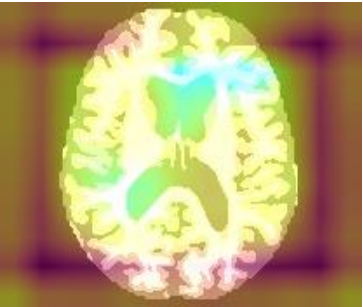


Fig. 16. Mapa de Calor resultante del experimento (ii)

W. S. Ucañay Barreto, and M. A. Coral Ygnacio, “Alzheimer's Diagnosis System based on Magnetic Resonance Imaging using the VGG16 Algorithm”, Latin-American Journal of Computing (LAJC), vol. 11, no. 1, 2024.

Para visualizar y resaltar los posibles cambios en la materia blanca del cerebro, se generó un mapa de calor con características similares al experimento anterior (Figura 16). En este mapa, se asignó un color celeste a las regiones donde se concentró el análisis y la atención del algoritmo, lo que indica la importancia de los cuencos cerebrales en la detección del Alzheimer leve. Además, se utilizaron los colores rojo y amarillo para resaltar áreas con niveles crecientes de intensidad o importancia en relación con los cambios estructurales detectados. Estas regiones resaltadas podrían ser indicativas de alteraciones específicas en la materia blanca asociadas con el Alzheimer en su etapa leve.

El uso del mapa de calor para la identificación de los cuencos cerebrales y la disminución de la materia blanca como enfoque principal en este experimento proporciona información relevante sobre las características distintivas de la enfermedad en su fase temprana. Estos hallazgos contribuyen a una mejor comprensión de los cambios estructurales asociados con el Alzheimer leve y pueden ser de utilidad en el diagnóstico y seguimiento de la enfermedad en etapas iniciales.

De los resultados anteriores, en la Figura 12 se puede observar que la clasificación obtenida para el experimento (i) fue de "Sin Alzheimer". Esta clasificación indica que, según los datos proporcionados por las imágenes de resonancia magnética (MRI), no se encontraron biomarcadores característicos del Alzheimer. Por lo tanto, se concluye que el sujeto no presenta signos de la enfermedad. En contraste, en el experimento (ii), la Figura 15 detalla los resultados, los cuales revelaron una clasificación de "Alzheimer Leve". Esto indica que, según los datos obtenidos de las imágenes de resonancia magnética (MRI), se encontraron indicios tempranos de la enfermedad de Alzheimer en el sujeto evaluado.

En ambos experimentos, se destaca que los resultados proporcionados por el sistema de diagnóstico del Alzheimer basado en imágenes pueden ser considerados como una herramienta complementaria para la toma de decisiones médicas. Sin embargo, en aplicaciones reales, es importante tener en cuenta que los resultados de los experimentos (i) y (ii) deben considerarse con cautela y no deben ser tomados como un diagnóstico concluyente.

De forma general, el sistema de diagnóstico basado en imágenes que se propone, demuestra ser una potencial herramienta útil para la identificación y clasificación de diferentes etapas del Alzheimer. No obstante, los resultados que arroja el sistema propuesto deben ser interpretados y validados por un médico especialista los cuales, luego de exámenes adicionales, posiblemente ayuden al diagnóstico definitivo de la enfermedad de Alzheimer en sus diferentes etapas.

A. Comparación con Trabajos Relacionados

Para evaluar el rendimiento de nuestro método, lo comparamos con otros métodos existentes para el diagnóstico de la enfermedad de Alzheimer (EA). La comparación se realizó utilizando diferentes conjuntos de datos, clases y tamaños de muestra, características, algoritmos y si tiene una clasificación binaria, múltiple o mixta.

TABLA 1. COMPARACIÓN CON MÉTODOS EXISTENTES PARA CLASIFICACIÓN BINARIA Y MÚLTIPLE

Nº	Autor	Base de Datos	Año	Binaria	Múltiple	% Precisión
1	[42]	ADNI	2023	Si	Si	77
2	[43]	ADNI	2022	Si	No	85.12
3	[44]	OASIS	2023	Si	No	95.48
4	[45]	ADNI	2023	Si	No	89
5	[46]	ADReSS	2021	Si	Si	84-90
6	[32]	ADNI	2021	Si	Si	80.9
7	[47]	ADNI	2021	Si	Si	81.73
8	[48]	OASIS	2021	Si	No	80
9	[49]	OASIS	2021	No	Si	72.8
10	[50]	OASIS	2022	No	Si	74.9
11	Propuesta	OASIS	2023	Si	Si	+82

La Tabla 1 presenta una comparativa de diversos estudios y las bases de datos empleadas, con el método propuesto ubicándose en un punto intermedio en términos de desempeño. Cuando se trata de la clasificación binaria (determinando si alguien presenta o no la enfermedad), el método propuesto se encuentra en línea con la media de los métodos existentes. Sin embargo, en el caso de la clasificación múltiple, donde se deben identificar distintas categorías relacionadas con la enfermedad, el método propuesto supera el rendimiento promedio. Cabe destacar que los estudios evaluados se basaban en una clasificación en cuatro categorías, al igual que el enfoque de este trabajo.

VI. CONCLUSIONES Y RECOMENDACIONES

A lo largo de este estudio, se demostró que el algoritmo de aprendizaje profundo VGG16, combinado con el análisis de imágenes de resonancia magnética, es un enfoque prometedor para el diagnóstico temprano de la enfermedad de Alzheimer. Se crearon módulos para el preprocesamiento de datos, entrenamiento y prueba del modelo, así como para la interfaz de usuario, asegurando la usabilidad del sistema. La metodología propuesta permitió alcanzar una precisión notable en la identificación de patrones de cambio estructural en el cerebro relacionados con esta enfermedad neurodegenerativa. Sin embargo, es importante destacar que, aunque la clasificación de una imagen de resonancia magnética es un componente crucial del diagnóstico, no es suficiente por sí sola. Los resultados de los experimentos han demostrado que el sistema es capaz de identificar y clasificar diferentes etapas de Alzheimer con una precisión razonable. Sin embargo, el diagnóstico final siempre debe ser realizado por profesionales médicos calificados, tomando en cuenta múltiples factores clínicos y pruebas adicionales.

En cuanto a los datos, a pesar de tener un conjunto de datos considerable de imágenes de resonancia magnética, la precisión del modelo podría haber sido mayor con un conjunto de datos más extenso y diverso. El desempeño del modelo

depende en gran medida de la cantidad y la calidad de los datos de entrenamiento. El entrenamiento del modelo de aprendizaje profundo y la implementación del sistema requirieron una capacidad de hardware significativa. No obstante, la infraestructura de hardware existente y las técnicas de optimización permitieron llevar a cabo el proyecto de manera eficiente.

A pesar de los resultados positivos obtenidos, se sugiere continuar con la investigación y el desarrollo en esta área para mejorar aún más la precisión y la eficacia de estas técnicas. Podría ser valioso explorar la integración de otros algoritmos de aprendizaje profundo o incluso enfoques híbridos, así como, el uso de librerías adicionales y actualizaciones que podrían mejorar la eficiencia del sistema.

Además, se recomienda implementar un proceso de validación más extenso que incluya una variedad más amplia de conjuntos de datos, es decir, recopilar más imágenes de resonancia magnética, preferiblemente de diversas etapas de la enfermedad y de diferentes poblaciones y grupos de edad. Asimismo, se sugiere que este sistema se utilice en conjunto con otras pruebas clínicas y de laboratorio. Todo esto para garantizar que el sistema pueda manejar una diversidad más amplia de casos.

Finalmente, para proyectos futuros de mayor escala, puede ser beneficioso invertir en hardware más potente o explorar soluciones en la nube para manejar el entrenamiento y la implementación de modelos de aprendizaje profundo. Es importante destacar que, aunque estas técnicas presentan un gran potencial, su uso debe complementar, y no reemplazar, las evaluaciones clínicas tradicionales llevadas a cabo por profesionales de la salud.

REFERENCIAS

[1] S. O. Danso, Z. Zeng, G. Muniz-Terrera, and C. W. Ritchie, “Developing an Explainable Machine Learning-Based Personalised Dementia Risk Prediction Model: A Transfer Learning Approach With Ensemble Learning Algorithms,” *Front. Big Data*, vol. 4, no. May, pp. 1–14, 2021, doi: 10.3389/fdata.2021.613047.

[2] G. Mukhtar and S. Farhan, “Convolutional neural network based prediction of conversion from mild cognitive impairment to alzheimer’s disease: A technique using hippocampus extracted from MRI,” *Adv. Electr. Comput. Eng.*, vol. 20, no. 2, pp. 113–122, 2020, doi: 10.4316/AECE.2020.02013.

[3] A. Gopalsamy, B. Radha, and K. Haridas, “Prediction of neurodegenerative disease using brain image analysis with multilinear principal component analysis and quadratic discriminant analysis,” *Int. J. Adv. Technol. Eng. Explor.*, vol. 9, no. 90, pp. 604–622, 2022, doi: 10.19101/IJATEE.2021.875325.

[4] M. Odusami, R. Maskeliūnas, R. Damaševičius, and S. Misra, “Explainable Deep-Learning-Based Diagnosis of Alzheimer’s Disease Using Multimodal Input Fusion of PET and MRI Images,” *J. Med. Biol. Eng.*, vol. 43, no. 3, pp. 291–302, 2023, doi: 10.1007/s40846-023-00801-3.

[5] T. O. Frizzell et al., “Artificial intelligence in brain MRI analysis of Alzheimer’s disease over the past 12 years: A systematic review,” *Ageing Res. Rev.*, vol. 77, no. March 2021, p. 101614, 2022, doi: 10.1016/j.arr.2022.101614.

[6] C. Kavitha, V. Mani, S. R. Srividhya, O. I. Khalaf, and C. A. Tavera Romero, “Early-Stage Alzheimer’s Disease Prediction Using Machine Learning Models,” *Front. Public Heal.*, vol. 10, no. March, pp. 1–13, 2022, doi: 10.3389/fpubh.2022.853294.

[7] M. EL-Geneedy, H. E. D. Moustafa, F. Khalifa, H. Khater, and E. Abdelhalim, “An MRI-based deep learning approach for accurate detection of Alzheimer’s disease,” *Alexandria Eng. J.*, vol. 63, pp. 211–221, 2023, doi: 10.1016/j.aej.2022.07.062.

[8] D. Agarwal, M. A. Berbis, T. Martín-Noguerol, A. Luna, S. C. P. Garcia, and I. de la Torre-Diez, “End-to-End Deep Learning Architectures Using 3D Neuroimaging Biomarkers for Early Alzheimer’s Diagnosis,” *Mathematics*, vol. 10, no. 15, pp. 1–28,

2022, doi: 10.3390/math10152575.

[9] M. Ansart et al., “Predicting the progression of mild cognitive impairment using machine learning: A systematic, quantitative and critical review,” *Med. Image Anal.*, vol. 67, p. 101848, 2021, doi: 10.1016/j.media.2020.101848.

[10] X. Song et al., “Graph convolution network with similarity awareness and adaptive calibration for disease-induced deterioration prediction,” *Med. Image Anal.*, vol. 69, p. 101947, 2021, doi: 10.1016/j.media.2020.101947.

[11] R. Jain, N. Jain, A. Aggarwal, and D. J. Hemanth, “Convolutional neural network based Alzheimer’s disease classification from magnetic resonance brain images,” *Cogn. Syst. Res.*, vol. 57, pp. 147–159, 2019, doi: 10.1016/j.cogsys.2018.12.015.

[12] M. Sudharsan and G. Thailambal, “Alzheimer’s disease prediction using machine learning techniques and principal component analysis (PCA),” *Mater. Today Proc.*, no. xxxx, 2021, doi: 10.1016/j.matpr.2021.03.061.

[13] T. Zhou, K. H. Thung, M. Liu, F. Shi, C. Zhang, and D. Shen, “Multi-modal latent space inducing ensemble SVM classifier for early dementia diagnosis with neuroimaging data,” *Med. Image Anal.*, vol. 60, 2020, doi: 10.1016/j.media.2019.101630.

[14] Z. Pei et al., “Alzheimer’s disease diagnosis based on long-range dependency mechanism using convolutional neural network,” *Multimed. Tools Appl.*, vol. 81, no. 25, pp. 36053–36068, 2022, doi: 10.1007/s11042-021-11279-z.

[15] M. Subramoniam, T. R. Apama, P. R. Anurenjan, and K. G. Sreeni, “Deep Learning-Based Prediction of Alzheimer’s Disease from Magnetic Resonance Images,” pp. 145–151, 2022, doi: 10.1007/978-981-16-7771-7_12.

[16] N. Mahendran and D. R. V. P. M., “A deep learning framework with an embedded-based feature selection approach for the early detection of the Alzheimer’s disease,” *Comput. Biol. Med.*, vol. 141, no. September 2021, p. 105056, 2022, doi: 10.1016/j.combiomed.2021.105056.

[17] G. Castellazzi et al., “A Machine Learning Approach for the Differential Diagnosis of Alzheimer and Vascular Dementia Fed by MRI Selected Features,” *Front. Neuroinform.*, vol. 14, no. June, pp. 1–13, 2020, doi: 10.3389/fninf.2020.00025.

[18] S. B. Çelebi and B. G. Emiroğlu, “A Novel Deep Dense Block-Based Model for Detecting Alzheimer’s Disease,” *Appl. Sci.*, vol. 13, no. 15, 2023, doi: 10.3390/app13158686.

[19] C. Y. Cheung et al., “A deep learning model for detection of Alzheimer’s disease based on retinal photographs: a retrospective, multicentre case-control study,” *Lancet Digit. Heal.*, vol. 4, no. 11, pp. e806–e815, 2022, doi: 10.1016/S2589-7500(22)00169-8.

[20] E. Lella, A. Pazienza, D. Lofù, R. Anglani, and F. Vitulano, “An ensemble learning approach based on diffusion tensor imaging measures for Alzheimer’s disease classification,” *Electron.*, vol. 10, no. 3, pp. 1–16, 2021, doi: 10.3390/electronics10030249.

[21] A. El-Gawady, M. A. Makhlof, B. S. Tawfik, and H. Nassar, “Machine Learning Framework for the Prediction of Alzheimer’s Disease Using Gene Expression Data Based on Efficient Gene Selection,” *Symmetry (Basel)*, vol. 14, no. 3, 2022, doi: 10.3390/sym14030491.

[22] J. Zhang et al., “Predicting future cognitive decline with hyperbolic stochastic coding,” *Med. Image Anal.*, vol. 70, p. 102009, 2021, doi: 10.1016/j.media.2021.102009.

[23] F. Al-khuzai, “PREDICTION OF ALZHEIMER’S DISEASE FROM 2-D ANATOMICAL MAGNETIC Institute of Graduate Studies Electrical and Computer Engineering PREDICTION OF ALZHEIMER’S DISEASE FROM 2-D ANATOMICAL MAGNETIC RESONANCE IMAGES USING DEEP LEARNING Fanar Emad Khazaal AL-,” no. November, 2022, doi: 10.13140/RG.2.2.13245.95205.

[24] A. Alsaedi, “On Prediction of Early Signs of Alzheimer’s—A Machine Learning Framework,” *ProQuest Diss. Theses*, p. 101, 2021, [Online]. Available: https://www.proquest.com/dissertations-theses/on-prediction-early-signs-alzheimer-s-machine/docview/2599058272/se-2?accountid=28534%0Ahttps://i-share-rsh.primo.exlibrisgroup.com/openurl/01CARLI_RSH/01CARLI_RS H:CARLI_RSH?url_ver=Z39.88-2004&rft_val_fmt=inf

[25] H. Li, M. Habes, D. A. Wolk, and Y. Fan, “A deep learning model for early prediction of Alzheimer’s disease dementia based on hippocampal magnetic resonance imaging data,” *Alzheimer’s Dement.*, vol. 15, no. 8, pp. 1059–1070, 2019, doi: 10.1016/j.jalz.2019.02.007.

W. S. Ucañay Barreto, and M. A. Coral Ygnacio, “Alzheimer's Diagnosis System based on Magnetic Resonance Imaging using the VGG16 Algorithm”, Latin-American Journal of Computing (LAJC), vol. 11, no. 1, 2024.

[26] C. Navarro, “Revisión de metodologías ágiles para el desarrollo de software.,” *Prospectiva*, vol. 11, no. 2, pp. 30–39, 2013, [Online]. Available: <http://www.redalyc.org/articulo.oa?id=496250736004%0ACómo>

[27] G. Rojas C., D. L. de Guevara, R. Jaimovich F., E. Brunetti, E. Faure L., and M. Gálvez M., “NEUROIMÁGENES EN DEMENCIAS,” *Rev. Médica Clínica Las Condes*, vol. 27, no. 3, pp. 338–356, 2016, doi: 10.1016/j.rmcl.2016.06.008.

[28] C. Wang et al., “A high-generalizability machine learning framework for predicting the progression of Alzheimer’s disease using limited data,” *npj Digit. Med.*, vol. 5, no. 1, pp. 1–10, 2022, doi: 10.1038/s41746-022-00577-x.

[29] S. Wu et al., “Application of artificial intelligence in clinical diagnosis and treatment: an overview of systematic reviews,” *Intell. Med.*, vol. 2, no. 2, pp. 88–96, 2022, doi: 10.1016/j.imed.2021.12.001.

[30] T. Habuza, N. Zaki, E. A. Mohamed, and Y. Statsenko, “Deviation from Model of Normal Aging in Alzheimer’s Disease: Application of Deep Learning to Structural MRI Data and Cognitive Tests,” *IEEE Access*, vol. 10, pp. 53234–53249, 2022, doi: 10.1109/ACCESS.2022.3174601.

[31] M. Liu, “Joint Classification and Regression via Deep Multi-Task Multi-Channel Learning for Alzheimer’s Disease Diagnosis,” *Physiol. Behav.*, vol. 176, no. 1, pp. 139–148, 2019, doi: 10.1109/TBME.2018.2869989.Joint.

[32] J. E. Arco, J. Ramírez, J. M. Górriz, and M. Ruz, “Data fusion based on Searchlight analysis for the prediction of Alzheimer’s disease,” *Expert Syst. Appl.*, vol. 185, 2021, doi: 10.1016/j.eswa.2021.115549.

[33] S. Basheer, S. Bhatia, and S. B. Sakri, “Computational Modeling of Dementia Prediction Using Deep Neural Network: Analysis on OASIS Dataset,” *IEEE Access*, vol. 9, pp. 42449–42462, 2021, doi: 10.1109/ACCESS.2021.3066213.

[34] A. V. Lebedev et al., “Random Forest ensembles for detection and prediction of Alzheimer’s disease with a good between-cohort robustness,” *NeuroImage Clin.*, vol. 6, pp. 115–125, 2014, doi: 10.1016/j.nicl.2014.08.023.

[35] C. V. Angkoso, H. P. A. Tjahyaningtijas, Y. Adrianto, A. D. Sensusiaty, I. K. E. Purnama, and M. H. Purnomo, “Multi-Features Fusion in Multi-plane MRI Images for Alzheimer’s Disease Classification,” *Int. J. Intell. Eng. Syst.*, vol. 15, no. 4, pp. 182–197, 2022, doi: 10.22266/ijies2022.0831.17.

[36] S. Toshkhujayev et al., “Classification of Alzheimer’s Disease and Mild Cognitive Impairment Based on Cortical and Subcortical Features from MRI T1 Brain Images Utilizing Four Different Types of Datasets,” *J. Healthe. Eng.*, vol. 2020, no. Mci, 2020, doi: 10.1155/2020/3743171.

[37] K. Shirbandi et al., “Accuracy of deep learning model-assisted amyloid positron emission tomography scan in predicting Alzheimer’s disease: A Systematic Review and meta-analysis,” *Informatics Med. Unlocked*, vol. 25, no. July, p. 100710, 2021, doi: 10.1016/j.imu.2021.100710.

[38] K. M. Poloni, I. A. Duarte de Oliveira, R. Tam, and R. J. Ferrari, “Brain MR image classification for Alzheimer’s disease diagnosis

using structural hippocampal asymmetrical attributes from directional 3-D log-Gabor filter responses,” *Neurocomputing*, vol. 419, pp. 126–135, 2021, doi: 10.1016/j.neucom.2020.07.102.

[39] M. K. Keles and U. Kilic, “Classification of Brain Volumetric Data to Determine Alzheimer’s Disease Using Artificial Bee Colony Algorithm as Feature Selector,” *IEEE Access*, vol. 10, no. July, pp. 82989–83001, 2022, doi: 10.1109/ACCESS.2022.3196649.

[40] Instituto Nacional de Ciencias Neurológicas Perú, “Instituto Nacional de Ciencias Neurológicas Perú.” <https://www.incn.gob.pe/> (accessed Oct. 29, 2023).

[41] D. S. Marcus, T. H. Wang, J. Parker, J. G. Csernansky, J. C. Morris, and R. L. Buckner, “Open Access Series of Imaging Studies (OASIS): Cross-sectional MRI Data in Young, Middle Aged, Nondemented, and Demented Older Adults,” *J. Cogn. Neurosci.*, vol. 19, no. 9, pp. 1498–1507, Sep. 2007, doi: 10.1162/jocn.2007.19.9.1498.

[42] P. Carcagni, M. Leo, M. Del Coco, C. Distant, and A. De Salve, “Convolution Neural Networks and Self-Attention Learners for Alzheimer Dementia Diagnosis from Brain MRI,” *Sensors*, vol. 23, no. 3, 2023, doi: 10.3390/s23031694.

[43] S. Liu et al., “Generalizable deep learning model for early Alzheimer’s disease detection from structural MRIs,” *Sci. Rep.*, vol. 12, no. 1, pp. 1–12, 2022, doi: 10.1038/s41598-022-20674-x.

[44] S. Lahmiri, “Integrating convolutional neural networks, kNN, and Bayesian optimization for efficient diagnosis of Alzheimer’s disease in magnetic resonance images,” *Biomed. Signal Process. Control*, vol. 80, p. 104375, Feb. 2023, doi: 10.1016/J.BSPC.2022.104375.

[45] X. Zheng, J. Cawood, C. Hayre, and S. Wang, “Computer assisted diagnosis of Alzheimer’s disease using statistical likelihood-ratio test,” *PLoS One*, vol. 18, no. 2 February, pp. 1–11, 2023, doi: 10.1371/journal.pone.0279574.

[46] E. Villatoro-Tello, S. P. Dubagunta, J. Fritsch, G. Ramírez-De-La-Rosa, P. Motlicek, and M. Magimai.-Doss, “Late fusion of the available lexicon and raw waveform-based acoustic modeling for depression and dementia recognition,” *Proc. Annu. Conf. Int. Speech Commun. Assoc. INTERSPEECH*, vol. 1, pp. 161–165, 2021, doi: 10.21437/Interspeech.2021-1288.

[47] S. D. Mishra and M. Dutta, “Modality feature fusion based Alzheimer’s disease prognosis,” *Optik (Stuttg.)*, vol. 272, p. 170347, 2023, doi: 10.1016/j.ijleo.2022.170347.

[48] C. L. Saratxaga et al., “Mri deep learning-based solution for alzheimer’s disease prediction,” *J. Pers. Med.*, vol. 11, no. 9, 2021, doi: 10.3390/jpm11090902.

[49] W. H. L. Pinaya et al., “Using normative modelling to detect disease progression in mild cognitive impairment and Alzheimer’s disease in a cross-sectional multi-cohort study,” *Sci. Rep.*, vol. 11, no. 1, pp. 1–13, 2021, doi: 10.1038/s41598-021-95098-0.

[50] S. Gupta, V. Saravanan, A. Choudhury, A. Alqahtani, M. R. Abonazel, and K. S. Babu, “Supervised Computer-Aided Diagnosis (CAD) Methods for Classifying Alzheimer’s Disease-Based Neurodegenerative Disorders,” *Comput. Math. Methods Med.*, vol. 2022, 2022, doi: 10.1155/2022/9092289.

AUTHORS



Werner Ucañay

Werner Shtaniklao Ucañay Barreto - Graduado en Ingeniería de Sistemas en el año 2023 de la Universidad Católica Sedes Sapientiae. Su línea de investigación apunta hacia el área de la Inteligencia Artificial y la Visión por Computadora. A lo largo de su formación académica, ha desarrollado habilidades y conocimientos en la aplicación de técnicas de Inteligencia Artificial para resolver problemas complejos. Su enfoque en la Visión por Computadora le ha permitido explorar y comprender cómo las máquinas pueden procesar y analizar imágenes y videos de manera eficiente, lo que tiene aplicaciones en una amplia gama de industrias, desde la salud hasta la automatización industrial.

Además, sus intereses incluyen la aplicación de la inteligencia artificial en el campo de la medicina, soluciones que aprovechen el poder de la inteligencia artificial para mejorar el diagnóstico médico, el análisis de imágenes médicas y la atención al paciente. Cree que la sinergia entre la tecnología y la medicina tiene el potencial de revolucionar la atención médica y mejorar la calidad de vida de las personas.



Marco Coral Ygnacio

Marco Antonio Coral Ygnacio es Magister en Ciencias en Ingeniería de Sistemas y Computación. Obtuvo su título de grado en la Universidad Inca Garcilaso de la Vega en 2001 y su maestría en ciencias en ingeniería de sistemas y computación en la misma universidad en 2006. Desde 2007, se desempeña como docente en la Facultad de Ingeniería de Sistemas e Informática de la Universidad Nacional Mayor de San Marcos y es docente investigador en la misma facultad. Además, desde 2011, forma parte del equipo docente de la Universidad Católica Sedes Sapientiae en Lima. Coral Ygnacio cuenta con una amplia experiencia en ingeniería de sistemas y computación, participando en proyectos relacionados con el desarrollo de sistemas de información, inteligencia artificial y aprendizaje automático. Además, ha llevado a cabo investigaciones centradas en la gestión administrativa en las universidades públicas peruanas. Su compromiso con la educación y la investigación se refleja en su carrera como docente e investigador en el campo de la ingeniería de sistemas y la gestión administrativa.

Exploring Topics in Information Technology Open Educational Resources through the LDA Algorithm

ARTICLE HISTORY

Received 18 September 2023
Accepted 27 November 2023

René Ludeña
Universidad Técnica Particular de Loja
Loja, Ecuador
rjludena@utpl.edu.ec

Verónica Segarra-Faggioni
Universidad Técnica Particular de Loja
Loja, Ecuador
Ecole de Technologie Supérieure, Quebec,
Canada
vasegarra@utpl.edu.ec
ORCID: 0000-0002-7275-7411

Audrey Romero-Peláez
Universidad Técnica Particular de Loja
Loja, Ecuador
Universidad Politécnica de Madrid Madrid,
España
aeromero2@utpl.edu.ec
ORCID: 0000-0002-3710-1257

Juan Carlos Morocho-Yunga
Universidad Técnica Particular de Loja
Loja, Ecuador
jcmorocho@utpl.edu.ec
ORCID: 0000-0001-6387-8054

Explorando Temas en Recursos Educativos Abiertos de Tecnologías de la Información a través del Algoritmo LDA

Exploring Topics in Information Technology Open Educational Resources through the LDA Algorithm

René Ludeña
Universidad Técnica Particular de Loja
Loja, Ecuador
rjludena@utpl.edu.ec

Verónica Segarra-Faggioni
Universidad Técnica Particular de Loja
Loja, Ecuador
Ecole de Technologie Supérieure,
Quebec, Canadá
vasegarra@utpl.edu.ec
ORCID: 0000-0002-7275-7411

Audrey Romero-Peláez
Universidad Técnica Particular de Loja
Loja, Ecuador
Universidad Politécnica de Madrid
Madrid, España
aeromero2@utpl.edu.ec
ORCID: 0000-0002-3710-1257

Juan Carlos Morochó-Yunga
Universidad Técnica Particular de Loja
Loja, Ecuador
jcmorochó@utpl.edu.ec
ORCID: 0000-0001-6387-8054

Abstract— This paper explores the application of machine learning and text mining techniques to discover OER issues in the context of Engineering Education. Applying the LDA (Latent Dirichlet Allocation) algorithm, themes are extracted from OER, it is possible to consider them as additional metadata. This augmentation serves to enhance the description and categorization of OER. Furthermore, this study introduces a methodology to automatically identify topics in open educational resources. In this research, a dataset of 80 OER was obtained from the Skills Commons repository. The highest coherence value achieved at 0.42, emerged when the number of topics was 9 in the LDA model. These nine topics are closely associated with Information Technology Education.

Keywords— metadata, OER, LDA, text mining, topic modeling

Resumen— Este artículo aplica el algoritmo Latent Dirichlet Allocation, LDA, como una técnica de aprendizaje de máquina y minería de texto para descubrir temas en OER en el contexto de la educación en ingeniería. El algoritmo LDA permite extraer temas, en este estudio los temas que se extraen de OER pueden ser considerados como metadatos adicionales que enriquecerán la descripción y clasificación de los mismos. Además, se define una metodología para la identificación automática de temas en los recursos educativos abiertos. En esta investigación, se utiliza un dataset de 80 OER extraído del repositorio Skills Commons. El valor más alto de coherencia es 0.42, cuando el número de temas en el modelo LDA es 9. Estos nueve temas están relacionados con Educación en Tecnologías de la Información.

Palabras clave— metadata, OER, LDA, minería de texto, modelado de temas

En la actualidad, la computación en la educación es un componente fundamental para el desarrollo de habilidades tecnológicas que impulsan el progreso de la sociedad. Con el avance constante de la inteligencia artificial (IA) y el aprendizaje automático (ML), es fundamental que las personas adquieran competencias en el campo de la informática para hacer frente a los desafíos de un mundo cada vez más digitalizado. En este contexto, los Recursos Educativos Abiertos (OER) han surgido como una solución prometedora para democratizar el acceso a materiales educativos de alta calidad.

Los Recursos Educativos Abiertos ofrecen una estrategia y oportunidad para mejorar el acceso a la educación de manera libre y abierta a través de sus materiales [1]. Los recursos educativos abiertos son materiales educativos de acceso gratuito para el público en general, los cuales proporcionan derechos de uso, reutilización y pueden adaptarse de acuerdo a las necesidades específicas. Con la abundancia de OER disponibles, descubrir contenido relevante y útil de manera eficiente se ha convertido en un desafío para los usuarios, debido a que los metadatos no siempre describen el recurso lo más completo posible [2].

Los metadatos son una colección de atributos que describen e identifican un recurso. En consecuencia, el uso de un único metadato estándar no es factible, ya que podría existir carencias en la cobertura de datos referentes a los recursos [3]. Debido a esta diversidad de información, los metadatos no siempre pueden controlar la calidad de los temas tratados en los recursos. Para abordar este tema, en este artículo exploramos el algoritmo LDA (Latent Dirichlet Allocation) como una técnica de aprendizaje de máquina y minería de texto para descubrir patrones y estructuras en los OER

R. Ludeña, V. Segarra-Faggioni, A. Romero-Pelaez, J.C. Morochó-Yunga
“Exploring Topics in Information Technology Open Educational Resources through the LDA Algorithm”,
Latin-American Journal of Computing (LAJC), vol. 11, no. 1, 2024.

seleccionados. El algoritmo LDA es una herramienta eficaz para analizar un conjunto de datos de texto y descubrir temas en estos documentos, por esta razón, se emplea LDA para identificar temas en recursos educativos abiertos de diferentes escenarios de formación en Ingeniería. Al emplear técnicas de aprendizaje automático no supervisado y minería de datos en los OER, se busca enriquecer la experiencia de aprendizaje al permitir que los usuarios encuentren los recursos apropiados según sus requerimientos de forma eficaz. Además, siguiendo la línea de investigación de varios autores que enfocan sus trabajos en la evaluación de recursos educativos abiertos desde un punto de vista técnico, se considera que con la aplicación de LDA se mejorará la organización y acceso al contenido educativo.

El presente trabajo está organizado de la siguiente manera: Sección 2. Se presentan trabajos relacionados al tema de estudio; Sección 3. Metodología aplicada; Sección 4. Experimentos con el algoritmo LDA, análisis y resultados obtenidos; y, Sección 5. Conclusiones y trabajos futuros.

II. TRABAJOS RELACIONADOS

El modelado de temas es una herramienta que brinda una visión de conjuntos de documentos individuales y las interconexiones entre ellos [4]. Varios autores han publicado y presentando resultados utilizando el algoritmo LDA en diferentes campos. En este trabajo, se realizó una búsqueda de artículos relacionados a la aplicación del algoritmo LDA en el contexto de los recursos educativos abiertos y se seleccionaron aquellos que aportan al objetivo de este trabajo (Ver Tabla I).

Una propuesta específica para analizar la calidad de los metadatos es aplicar el algoritmo Random Forest (RF) [5], el análisis está enfocado en aplicar un modelo de predicción para anticipar la calidad de los OER. Mediante el uso de metadatos como característica de entrada, los hallazgos indican que este modelo alcanza el 94.6% de precisión al identificar de manera correcta los OER de alta calidad.

Por otro lado, el estudio de [6] combina la técnica *ElasticSearch* y el modelo LDA con la finalidad de clasificar artículos académicos de acuerdo a temas relevantes y características extraídas de los resúmenes de los documentos. Se utilizaron como variables de entrada las palabras clave y el resumen de los artículos académicos.

Otro estudio [7] compara los métodos: LDA, BERT y Tf-idf para determinar el más eficaz para la clasificación de texto en base a una etiqueta predefinida. El conjunto de datos en este estudio corresponde a temas deportivos y educativos que fueron recopilados por estudiantes de postgrado. Así también, el estudio de [8] extrae temas de OER, aplicando técnicas de minería de texto, para generar metadatos de alta calidad.

Los trabajos relacionados utilizan Random Forest y BERT para tareas de clasificación y regresión utilizando etiquetas en un documento, mientras que LDA permite extraer temas de documentos. El enfoque de utilizar el algoritmo LDA para investigar el tema de documentos o palabras en diversos campos porque permite analizar grandes cantidades de datos. En este estudio, se aplica LDA para descubrir temas dentro de los recursos educativos abiertos con la finalidad de

considerarlos como metadatos adicionales que enriquecerán la descripción y clasificación de los OER.

TABLA I. RESUMEN DE TRABAJOS RELACIONADOS

Propósito	Metodología	Resultados
Explorar la relación entre la calidad de los metadatos y la calidad de OER [5]	1. Análisis de datos exploratorio sobre los metadatos de OER de Youtube, aplicando RF. 2. Predicción basada en metadatos para anticipar la calidad de los OER.	La clasificación de OER de alta calidad alcanza una precisión del 94,6%.
Clasificar automáticamente artículos académicos [6]	1. Aplica el método de clasificación ElasticSearch y el modelo LDA basado en temas para extraer las características de los artículos académicos. La similitud semántica utiliza palabras clave como las variables de entrada.	El valor de $k = 50$ utilizado para aplicar LDA.
Explorar la aplicación práctica de tres métodos de modelización (LDA, BERT y TF-IDF) y determinar el método más eficaz para la clasificación de documentos. [7]	1.Preprocesamiento de los datos. 2.1. Calcular Tf – idf (Frecuencia de término – Documento inverso Frecuencia) 2.2. Aplicar algoritmo LDA: se crea un diccionario y una bolsa de palabras. 2.3. Aplicación del método BERT: se crea vectores de incrustación de oraciones.	El método BERT alcanzó el 92,6% de éxito, en la clasificación de documentos.
Realizar la extracción de temas de OER mediante técnicas de minería de texto para generar metadatos de alta calidad, lo que puede ayudar a los alumnos a construir itinerarios de aprendizaje eficaces hacia sus objetivos de aprendizaje individuales [8].	1.Recolección de datos 2.Preprocesamiento de datos 3. Aplicación del algoritmo Latent Dirichlet Allocation (LDA) 3.1.Calcular la coherencia para seleccionar el valor de k 4. Evaluación del modelo utilizando la métrica $F1-score$	El modelo extrajo temas de OER con un 79% de puntuación F1.

III. METODOLOGÍA

En el presente trabajo, se aplica el algoritmo LDA (Latent Dirichlet Allocation) para la identificación automática de temas que se aborden en los recursos educativos abiertos utilizados en la formación en Ingeniería.

Con este objetivo se aplicaron las fases: entender el proceso, entender los datos, preparar los datos, construir el modelo y validar el modelo. La Figura 1 presenta el esquema de los pasos propuestos.

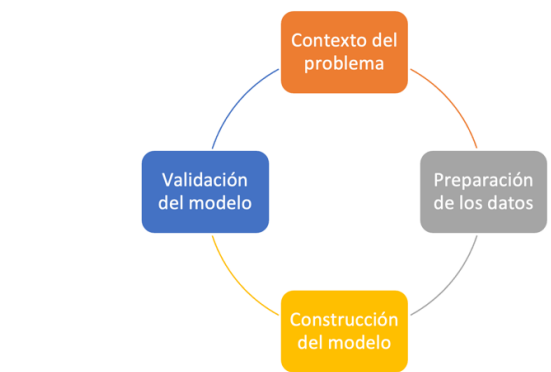


Fig. 1. Metodología aplicada

A. Contexto del problema

Se realiza una definición técnica del problema y, además, precisa los aspectos del dominio de trabajo, en este caso los metadatos de los repositorios OER. Por lo tanto, en esta fase inicial se busca comprender el problema de la búsqueda y recomendación eficiente de los recursos educativos abiertos.

En los repositorios de OER, los metadatos de estos recursos no brindan la eficiencia necesaria para los servicios de búsqueda [11]. Por esta razón, se plantea aplicar el algoritmo LDA, como técnica de aprendizaje de máquina y minería de datos, que permita la extracción de los temas que son tratados en un OER y de esta manera proveer una alternativa que mejore la eficiencia de los usuarios al encontrar y seleccionar materiales según sus requerimientos.

El algoritmo LDA permite identificar la distribución del documento sobre los temas y la distribución de un tema en función de las palabras observadas [12]. De esta manera, se identificarán automáticamente temas de los OER, los mismos que pueden ser considerados como metadatos adicionales que enriquecerán la descripción y clasificación de los recursos educativos abiertos seleccionados.

Además, se realiza un análisis exploratorio con el fin de obtener un panorama de lo que se puede conseguir a través de los datos existentes. En este punto, el conocimiento del dominio de trabajo permite guiar este análisis. Así también, en esta fase se identifican posibles problemas de calidad en los datos y se detectan subconjuntos en los datos que podrían ser interesantes [12].

B. Preparación de los datos

La fase de preparación de datos lleva a cabo todas las actividades para construir un conjunto de datos adecuado y listo para utilizar en modelos de aprendizaje automático. En esta fase el propósito es obtener una comprensión integral del conjunto de documentos que se obtengan del conjunto de datos (dataset). Los metadatos seleccionados del conjunto de datos son los siguientes: *id*, *título*, *url* y *type* (*tipo de material principal*). Con base en estos metadatos, se lleva a cabo la recopilación, exploración y evaluación de los datos para asegurar que se cuente con una colección representativa de recursos educativos abiertos y se abarque una amplia variedad de temáticas utilizadas en la formación en ingeniería.

Durante esta fase, los documentos recopilados pasan por un proceso de preprocesamiento para asegurar que sean adecuados para el análisis [13]. Esta etapa implica diversas

tareas de procesamiento de lenguaje natural (en inglés NLP) como: limpieza de texto para eliminar símbolos irrelevantes, puntuación y caracteres especiales y la eliminación de palabras vacías que no aportan información relevante para el análisis. Además, se realiza la *tokenización* para identificar las palabras significativas y la lematización para reducir las palabras a su forma raíz y evitar redundancias con el fin de mejorar la eficacia del análisis.

C. Construcción del modelo

Durante el proceso de construcción del modelo se eligen y aplican diversas técnicas de modelado, mientras se ajustan los parámetros para alcanzar valores óptimos. En muchos casos, se dispone de múltiples técnicas para abordar el mismo tipo de problema. Algunas técnicas demandan formatos de datos particulares, resaltando así la relación entre la preparación de los datos y el proceso de modelado [9].

En esta fase, se incorpora el algoritmo LDA para extraer temas de los documentos que han sido previamente preprocesados. LDA identifica estos temas basándose en la distribución probabilística de palabras en los documentos y la distribución de temas dentro de los mismos [14]. Es importante determinar el número adecuado de temas, lo que dependerá del conocimiento del dominio y el nivel de granularidad deseado. Una vez definido el modelo basado en el algoritmo de aprendizaje LDA se ejecuta el entrenamiento y validación del modelo.

D. Validación del modelo

En esta etapa final, luego de contar con uno o más modelos que pueden tener alta calidad, desde una perspectiva de análisis de datos. Antes de avanzar con la implementación final del modelo, es necesario realizar una evaluación exhaustiva y revisar los pasos ejecutados para la construcción del modelo, con la finalidad de asegurarse que logre los objetivos propuestos para solucionar el problema planteado al inicio. Al concluir esta fase, se debe llevar a cabo un análisis detallado de los resultados [9], [15]. Los resultados de esta etapa son fundamentales para ajustar y mejorar el modelo LDA y, en última instancia, para lograr un análisis más preciso y útil.

La fase de validación permite garantizar que los temas extraídos sean de alta calidad y puedan proporcionar una comprensión significativa de los contenidos en los OER. La evaluación de la calidad y coherencia de los temas extraídos por el modelo LDA se realiza mediante métricas específicas que miden la relación semántica entre las palabras dentro de cada tema.

IV. ANÁLISIS Y DISCUSIÓN DE RESULTADOS

A. Conjunto de datos

El conjunto de datos se recopiló del repositorio Skills Commons, una plataforma que cuenta con una amplia variedad de Recursos Educativos Abiertos (OER). Este conjunto de datos incluye metadatos como: *título*, *url*, *type* (tipo de material principal), entre otros. Los OER utilizados en este trabajo corresponden al tema “*Tecnologías de la Información*” y los metadatos y recursos se encuentran en idioma inglés.

Para la construcción del corpus, se seleccionó el metadato *type* que significa “tipo de material principal” ya que permite identificar el tipo de recurso educativo abierto. En este estudio,

se seleccionaron aquellos recursos que están etiquetados como “*Final Program Report*” en el metadato *type* (Ver Figura 2.). Este criterio proporcionó un total de 80 OER que son de interés para el propósito de este trabajo de investigación. Cada uno de estos recursos contiene una colección de 103 documentos relacionados con diversos aspectos de las *Tecnologías de la Información*.

```
dataFt= data.loc[data['type'] == 'Final Program Report']
dataOp= dataFt.loc[:, ['title','url', 'type']]
dataOpc= dataOp
```

	title	url	type
id			
13172	RITA Consortium Final Evaluation Report Septem...	https://www.skillscommons.org/handle/faaccct/...	Final Program Report
13748	Third party evaluation	https://www.skillscommons.org/handle/faaccct/...	Final Program Report
18125	Get IT Project Evaluation Final Report	https://www.skillscommons.org/handle/faaccct/...	Final Program Report
18467	Third Party Evaluations of the IT Programs	https://www.skillscommons.org/handle/faaccct/...	Final Program Report
18557	Final Evaluation Report: Health Information Te...	https://www.skillscommons.org/handle/faaccct/...	Final Program Report
15621	New River Community and Technical College's Fr...	https://www.skillscommons.org/handle/faaccct/...	Final Program Report
15677	Final Evaluation Report Developing Pathways fo...	https://www.skillscommons.org/handle/faaccct/...	Final Program Report
18509	New Jersey Health Professions Pathways to Regi...	https://www.skillscommons.org/handle/faaccct/...	Final Program Report
15562	FINAL EXTERNAL EVALUATION REPORT - Clovis Com...	https://www.skillscommons.org/handle/faaccct/...	Final Program Report
9295	RCC TAACCCT Final Evaluation Report	https://www.skillscommons.org/handle/faaccct/...	Final Program Report

Fig. 2. Ejemplo de recursos educativos abiertos etiquetados como “Final Program Report”

B. Preparación de los datos

Se recopilaron los documentos relacionados con diversos aspectos de las *Tecnologías de la Información* para el análisis y construcción del modelo. Para la preparación y el preprocesamiento de datos se utilizó el lenguaje de programación Python, reconocido por su amplia gama de librerías especializadas en el procesamiento de lenguaje natural, mediante la biblioteca *nlTK* [16]. A continuación se realizan: la implementación del algoritmo LDA, la visualización de datos y la evaluación del modelo [17].

Los pasos realizados en esta fase de preparación y preprocesamiento de datos se detallan en la Tabla II.

C. Representación de documentos (Tf-idf)

Una vez procesados los datos, se realizó el análisis de frecuencia de palabras considerando que es común encontrar palabras que aparecen en múltiples documentos, sin importar la colección de datos a la que pertenecen. La técnica Tf-idf es ampliamente utilizada para reducir el peso de las palabras que aparecen con frecuencia en los vectores de características [14]. Esta técnica funciona mediante la asignación de una ponderación a los términos según la frecuencia de aparición en los documentos, lo que ayuda a destacar las palabras clave que realmente aportan significado a los documentos.

Para aplicar el algoritmo LDA a un conjunto de documentos, primero LDA asume un número fijo de temas (temas), y las relaciones *tema-palabra* y *tema-documento* (*bag of words*) se modelan con matrices de probabilidades *palabra-en-tema* y *tema-en-documento* (Tf-idf).

D. Aplicación del algoritmo LDA

Se utilizó la librería *gensim* en Python con el propósito de construir y entrenar el modelo LDA. En este proceso, los parámetros “*alpha*” y “*eta*” se configuraron con valores automáticos. Estos parámetros afectan la densidad de distribución de temas de los documentos y la palabra perteneciente a un tema.

TABLA II. PREPARACIÓN Y PREPROCESAMIENTO DE DATOS

	Tarea	Descripción
Preparación de datos	Descarga de documentos por OER	Se descargaron los documentos que forman parte del OER y se almacenó en carpetas identificadas por el “ <i>id</i> ”.
	Formateo de documentos	Se unificó el formato de los documentos, aquellos que estaban en formato PDF se convirtieron en formato DOCX.
	Transformación de documentos	Los metadatos de los 80 OER seleccionados y el contenido de los documentos de cada OER se almacenó en un archivo CSV para facilitar el manejo y tratamiento de los datos.
Preprocesamiento	Normalización	Para reducir el ruido se transforman las palabras a minúscula, se eliminan siglas, signos de puntuación, números y caracteres especiales.
	Tokenización	Se divide el texto en unidades más pequeñas llamadas tokens. Para tokenizar el texto se utiliza la función <i>nlTK.word.tokenize()</i> .
	Eliminación palabras vacías	Se eliminan palabras muy comunes, por ejemplo, “ <i>el</i> ”, “ <i>la</i> ”, “ <i>y</i> ”, “ <i>de</i> ”, que aparecen con frecuencia en el texto pero que no aportan un significado importante.
	Identificación de bigramas y trigramas	Se identifican y agrupan las palabras adyacentes en dos y tres elementos consecutivos, respectivamente, para facilitar la representación adecuada del texto.
	Lematización	Se reducen las palabras a su forma base o raíz, conocida como lema. Para la lematización se utilizó la librería <i>WordNetLemmatizer</i> del paquete <i>nlTK</i> [16]. En este estudio, la lematización se realiza considerando que los tokens son verbos.

E. Validación y resultados

En esta sección, para evaluar el modelo generado por el algoritmo LDA, se aplica la medida de coherencia de temas (*c_v*) para analizar el desempeño de un conjunto de modelos de temas de diferentes tamaños y tipos. La puntuación de coherencia LDA es una métrica utilizada para evaluar la calidad de los grupos de palabras generados por el modelo LDA. La puntuación de coherencia más alta proporciona una medida cuantitativa que muestra el grado de relación entre las palabras y un tema.

En la figura 3, se presentan los resultados de los valores de coherencia obtenidos para diferentes números de temas (*k*) al

aplicar el modelo de matriz Tf-idf con bigramas y trigramas en un análisis de temas. Los valores de k van desde 2 hasta 12 y se selecciona el valor más alto. Se puede observar que el valor más alto de coherencia (0.42) se obtiene cuando el valor óptimo de $k = 9$. Esto indica que el modelo de 9 temas tiene una coherencia de temas relativamente alta, lo que sugiere que las palabras dentro de los temas están bien relacionadas y representan patrones claros en el corpus.

Cabe recalcar que la medida de validación aplicada es coherencia c_v , la que es utilizada en el contexto de modelado de temas que permite medir la coherencia semántica entre palabras en un modelo LDA. Mientras que la métrica F1-score utilizada en [8] se aplica para la evaluación de modelos de clasificación.

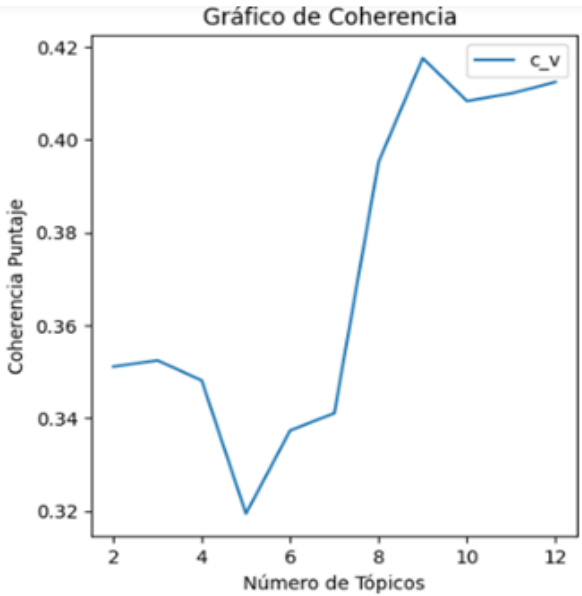


Fig. 3. Métrica de coherencia

Finalmente, para tener una visualización interactiva de los temas y palabras clave se utilizó la librería *pyLDAvis*. Esta librería crea el mapa de distancias entre temas y su distribución de temas para cada OER sobre *Tecnologías de la Información*, con el fin de ayudar en la interpretación de los temas en el modelado (Ver Figura 4).

En la figura 4, se presenta el panel derecho del mapa de distancia, se visualizan las 30 palabras más frecuentes de cada tema calculado con el modelo LDA. Cada tema se representa con una lista de palabras ponderadas por sus respectivos valores TF-IDF. Mientras que, en la interpretación de temas, se analizan las palabras más probables asociadas con cada tema para entender su significado y relevancia.

En la figura 5, se presenta el panel izquierdo del mapa de distancias donde se visualiza el tamaño de cada tema que representa la prevalencia sobre el conjunto de datos. Las burbujas más grandes indican que el tema es más prevalente en comparación con las burbujas más pequeñas.

Los temas 4 y 7 tienen una naturaleza independiente de otros temas. El tema 9 se integra con los temas 1 y 8, es decir, cuentan o comparten términos similares. Entre los términos más frecuentes están “program”, “student”, “college”, “participant” y “community” que están relacionados con Educación en Tecnologías de la Información.

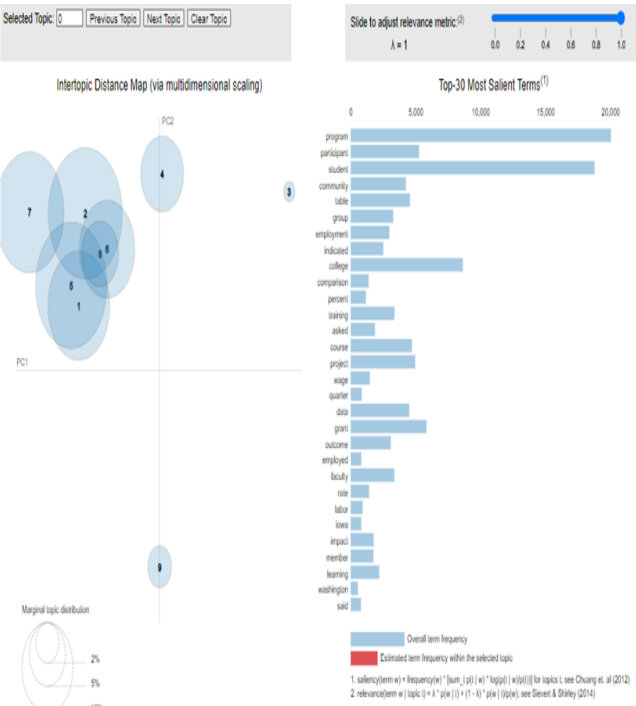


Fig. 4. Visualización de las 30 palabras más frecuentes

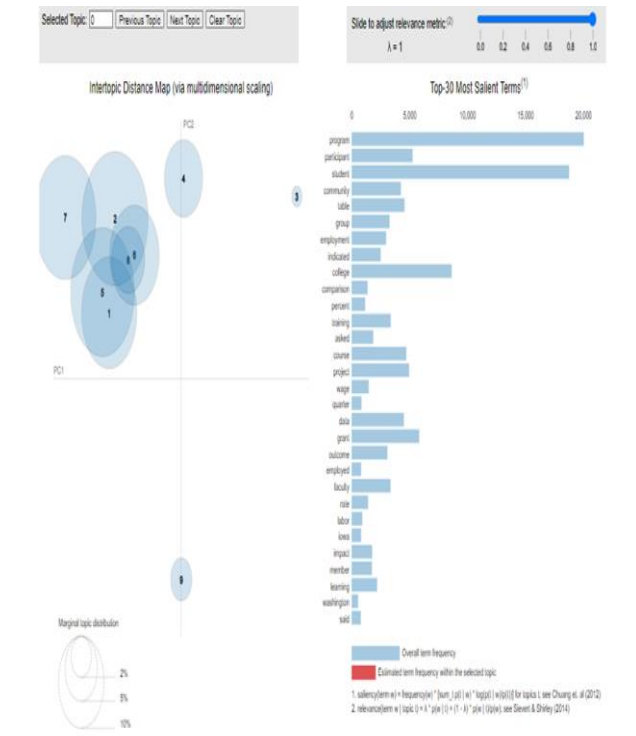


Fig. 5. Mapa de distancia de temas en OER sobre "Tecnologías de la información"

V. CONCLUSIONES Y TRABAJO FUTURO

El modelado de temas cumple un papel importante en la informática para la minería de textos. En este estudio, se aplicó el algoritmo LDA a un conjunto de documentos de Recursos Educativos Abiertos (OER) con el fin de identificar y analizar los temas presentes en toda la colección, así como en cada documento individual y las relaciones entre los documentos. Esta técnica es popular y relativamente simple de

implementar, sin embargo, puede ser sensible a la elección de un número de temas específicos.

En este artículo, se presentan los resultados de los experimentos y análisis realizados de OER extraídos del repositorio Skill Commons mediante la aplicación del algoritmo LDA. Los resultados muestran que la mayoría de los temas abordados están relacionados con Educación en Tecnologías de la Información. Estos resultados proporcionan al usuario una forma de identificar los temas principales en los OER sin necesidad de realizar una revisión detallada de cada recurso.

Se propone en futuros experimentos incluir métodos de aprendizaje profundo que permitan incluir más factores para construir un modelo de predicción orientado a la comparación de metadatos extrayendo las características de los OER.

AGRADECIMIENTO

El equipo de investigadores agradece a la Universidad Técnica Particular de Loja, especialmente al Grupo de Investigación de Sistemas Basados en Conocimiento (KBS-RG); y, a la SENESCYT de Ecuador.

REFERENCIAS

[1] V. Segarra-Faggioni and A. Romero-Pelaez, “Automatic classification of OER for metadata quality assessment,” in 2022 International Conference on Advanced Learning Technologies (ICALT), 2022, pp. 16–18. doi: 10.1109/ICALT55010.2022.00011.

[2] X. Ochoa and E. Duval, “Automatic evaluation of metadata quality in digital repositories,” Int. J. Digit. Libr., vol. 10, no. 2–3, pp. 67–91, Aug. 2009, doi: 10.1007/s00799-009-0054-4.

[3] J. Chicaiza, N. Piedra, J. Lopez-Vargas, and E. Tovar-Caro, “Recommendation of open educational resources. An approach based on linked open data,” null, 2017, doi: 10.1109/educon.2017.7943018.

[4] H. Jelodar et al., “Latent Dirichlet allocation (LDA) and topic modeling: models, applications, a survey,” Multimed. Tools Appl., vol. 78, no. 11, pp. 15169–15211, Jun. 2019, doi: 10.1007/S11042-018-6894-4/TABLES/11.

[5] M. Tavakoli, M. Elias, G. Kismihók, and S. Auer, “Metadata analysis of open educational resources,” in ACM International Conference Proceeding Series, 2021, pp. 626–631. doi: 10.1145/3448139.3448208.

[6] M. Kim and D. Kim, “A Suggestion on the LDA-Based Topic Modeling Technique Based on Elasticsearch for Indexing Academic

Research Results,” Appl. Sci., vol. 12, no. 6, p. 3118, Mar. 2022, doi: 10.3390/app12063118.

[7] S. Ozdemirci and M. Turan, “Case Study on well-known Topic Modeling Methods for Document Classification,” in 2021 6th International Conference on Inventive Computation Technologies (ICICT), Jan. 2021, pp. 1304–1309. doi: 10.1109/ICICT50816.2021.9358473.

[8] M. Molavi, M. Tavakoli, and G. Kismihók, “Extracting Topics from Open Educational Resources,” in Addressing Global Challenges and Quality Education, 2020, pp. 455–460.

[9] R. Wirth and J. Hipp, “Crisp-dm: towards a standard process modell for data mining,” 2000. [Online]. Available: <https://api.semanticscholar.org/CorpusID:1211505>

[10] P. Haya, “La metodología CRISP-DM en ciencia de datos,” INSTITUTO DE INGENIERÍA DEL CONOCIMIENTO, 2021. <https://www.iic.uam.es/innovacion/metodologia-crisp-dm-ciencia-de-datos/>

[11] M. Tavakoli, M. Elias, G. Kismihok, and S. Auer, “Quality Prediction of Open Educational Resources A Metadata-based Approach,” in 2020 IEEE 20th International Conference on Advanced Learning Technologies (ICALT), Jul. 2020, pp. 29–31. doi: 10.1109/ICALT49669.2020.00007.

[12] D. M. Blei, A. Y. Ng, and M. Jordan, “Latent Dirichlet Allocation,” J. Mach. Learn. Res., vol. 3, pp. 993–1022, 2003, Accessed: Mar. 25, 2019. [Online]. Available: <http://www.jmlr.org/papers/volume3/blei03a/blei03a.pdf>

[13] H. Lane, C. Howard, and H. Max Hapke, Natural Language Processing in Action: Understanding, Analyzing, and Generating Text with Python. Manning Publications., 2019.

[14] V. Mirjalili and S. Raschka, Python Machine Learning. Marcombo, 2019.

[15] S. Weiss, N. Indurkha, and T. Zhang, Fundamentals of Predictive Text Mining, Second Edi. 2015. doi: 10.1007/978-1-4471-6750-1.

[16] S. Bird and E. Loper, “NLTK: The Natural Language Toolkit,” 2006. Accessed: Oct. 14, 2018. [Online]. Available: www.python.org.

[17] S. Raschka, J. Patterson, and C. Nolet, “Machine Learning in Python: Main developments and technology trends in data science, machine learning, and artificial intelligence,” CoRR, vol. abs/2002.0, 2020, [Online]. Available: <https://arxiv.org/abs/2002.04803>

AUTHORS



René Ludeña

Estudiante de la carrera de Sistemas informáticos y Computación, Universidad Técnica de Loja.
Participó en el proyecto "Sistema de automatización de riesgos operativos "SARO" del Banco de Loja" donde se implementaron diferentes tecnologías para optimizar procesos, mejorar la eficiencia y prevenir los riesgos, 2021.



Verónica Segarra

Ph.D en Ingeniería por la École de technologie supérieure, Montreal -Canadá, 2021. La tesis de doctorado fue realizada bajo la dirección de la Prof. Sylvie Ratté, Ph.D. de la ÉTS, Montréal, Canadá y codirección del Prof. Frank de Jong de la Universidad AERES, Wageningen, Países Bajos.
Participó en un Hackaton en Países Bajos donde se evaluó los informes escritos de los alumnos de maestría de la Universidad AERES, 2018.
Realizó pasantías de investigación en CRIM (Centre de Recherche Informatique de Montréal) como parte del equipo "Parole et Texte" para análisis de texto utilizando técnicas de procesamiento de lenguaje natural, 2019.
Magíster en Auditoría de Gestión de la Calidad, Universidad Técnica Particular de Loja -Ecuador, 2008.
Diploma en Mejora de Procesos, Tecnológico de Monterrey, 2012.
Ingeniera en Sistemas Informáticos y Computación, Universidad Técnica Particular de Loja - Ecuador.
Licenciada en Ciencias de la Educación Mención Inglés, Universidad Técnica Particular de Loja - Ecuador.



Audrey Romero

Profesor titular a tiempo completo del Departamento de Ciencias de la Computación y Electrónica en la Universidad Técnica de Loja.
Máster en Ciencias y Tecnologías de la Computación por la Universidad Politécnica de Madrid.
Sus campos de investigación son: Educación abierta, Modelos de Calidad, Tecnologías Semánticas, Análisis de Datos y Tecnologías Emergentes.



Juan Carlos Morocho

Profesor titular a tiempo completo del Departamento de Ciencias de la Computación y Electrónica en la Universidad Técnica de Loja.
Graduado de la Maestría de Inteligencia Artificial en la Universidad Politécnica de Madrid.
Sus campos de investigación son: Diseño Ontológico, Tecnologías Semánticas y Análisis de Datos.

LAJC

LATIN-AMERICAN JOURNAL OF COMPUTING

Published by

Escuela Politécnica Nacional
Facultad de Ingeniería de Sistemas

Quito-Ecuador

<https://lajc.epn.edu.ec/>
lajc@epn.edu.ec

January 2024



LJJC

Vol XI, Issue 1, January 2024

LAJC
LATIN-AMERICAN
JOURNAL OF
COMPUTING