

EDUCACIÓN

SECRETARÍA DE EDUCACIÓN PÚBLICA



Instituto Politécnico Nacional
"La Técnica al Servicio de la Patria"

Research in Computing Science

Vol. 154 No. 6
June 2025

Research in Computing Science

Series Editorial Board

Editors-in-Chief:

Grigori Sidorov, CIC-IPN, Mexico
Gerhard X. Ritter, University of Florida, USA
Jean Serra, Ecole des Mines de Paris, France
Ulises Cortés, UPC, Barcelona, Spain

Associate Editors:

Jesús Angulo, Ecole des Mines de Paris, France
Jihad El-Sana, Ben-Gurion Univ. of the Negev, Israel
Alexander Gelbukh, CIC-IPN, Mexico
Ioannis Kakadiaris, University of Houston, USA
Petros Maragos, Nat. Tech. Univ. of Athens, Greece
Julian Padget, University of Bath, UK
Mateo Valero, UPC, Barcelona, Spain
Olga Kolesnikova, ESCOM-IPN, Mexico
Rafael Guzmán, Univ. of Guanajuato, Mexico
Juan Manuel Torres Moreno, U. of Avignon, France
Miguel González-Mendoza, ITESM, Mexico

Editorial Coordination:

Alejandra Ramos Porras

RESEARCH IN COMPUTING SCIENCE, Año 25, Volumen 154, No. 6, Junio de 2025, es una publicación mensual, editada por el Instituto Politécnico Nacional, a través del Centro de Investigación en Computación. Av. Juan de Dios Bátiz S/N, Esq. Av. Miguel Othón de Mendizábal, Col. Nueva Industrial Vallejo, C.P. 07738, Ciudad de México, Tel. 57 29 60 00, ext. 56571. <https://www.rcs.cic.ipn.mx>. Editor responsable: Dr. Grigori Sidorov. Reserva de Derechos al Uso Exclusivo del Título No. 04-2019-082310242100-203. ISSN: en trámite, otorgado por el Instituto Nacional del Derecho de Autor. Responsable de la última actualización de este número: el Centro de Investigación en Computación, Dr. Grigori Sidorov, Av. Juan de Dios Bátiz S/N, Esq. Av. Miguel Othón de Mendizábal, Col. Nueva Industrial Vallejo, C.P. 07738. Fecha de última modificación 06 de Junio de 2025.

RESEARCH IN COMPUTING SCIENCE, Year 25, Volume 154, No. 6, June, 2025, is a monthly publication edited by the National Polytechnic Institute through the Center for Computing Research. Av. Juan de Dios Bátiz S/N, Esq. Miguel Othón de Mendizábal, Nueva Industrial Vallejo, C.P. 07738, Mexico City, Tel. 57 29 60 00, ext. 56571. <https://www.rcs.cic.ipn.mx>. Editor in charge: Dr. Grigori Sidorov. Reservation of Exclusive Use Rights of Title No. 04-2019-082310242100-203. ISSN: pending, granted by the National Copyright Institute. Responsible for the latest update of this issue: the Computer Research Center, Dr. Grigori Sidorov, Av. Juan de Dios Bátiz S/N, Esq. Av. Miguel Othón de Mendizábal, Col. Nueva Industrial Vallejo, C.P. 07738. Last modified on June 6, 2025.

Advances in Artificial Intelligence

Bella Citlali Martínez-Seis (ed.)



Instituto Politécnico Nacional, Centro de Investigación en Computación
México 2025

ISSN: in process

Copyright © Instituto Politécnico Nacional 2025
Formerly ISSN: 1870-4069, 1665-9899

Instituto Politécnico Nacional (IPN)
Centro de Investigación en Computación (CIC)
Av. Juan de Dios Bátiz s/n esq. M. Othón de Mendizábal
Unidad Profesional “Adolfo López Mateos”, Zacatenco
07738, México D.F., México

<http://www.rcs.cic.ipn.mx>
<http://www.ipn.mx>
<http://www.cic.ipn.mx>

The editors and the publisher of this journal have made their best effort in preparing this special issue, but make no warranty of any kind, expressed or implied, with regard to the information contained in this volume.

All rights reserved. No part of this publication may be reproduced, stored on a retrieval system or transmitted, in any form or by any means, including electronic, mechanical, photocopying, recording, or otherwise, without prior permission of the Instituto Politécnico Nacional, except for personal or classroom use provided that copies bear the full citation notice provided on the first page of each paper.

Indexed in LATINDEX, DBLP and Periodica

Electronic edition

Table of Contents

	Page
Early Corruption Detection System Using Machine Learning Algorithms 5 <i>Maria Fernanda Heredia-Soto, Edgar Huaranga, Franci Suni-Lopez</i>	5
Metodología BRAIN: Una metodología para la construcción de chatbots basados en BERT y sus aplicaciones en empresas 13 <i>Miguel Angel González Reyes, Francisco Edgar Castillo Barrera, Héctor Gerardo Pérez González</i>	13
Transformer Architecture Reconfiguration Using Heterogeneous Attention 21 <i>Xamanek Martínez-Marín, José Adan Hernández-Nolasco, Noel Zacarias Morales</i>	21
Una aproximación al marco jurídico de la Inteligencia Artificial en México 35 <i>Antonio de Jesús Remes Díaz</i>	35
Evaluación preliminar de reconocimiento de patrones en gotas secas de metotrexato 47 <i>Carlos A. Martínez-Miwa, R. Magdalena Sánchez-Albores, Yohana J. P. Carreón, Jorge Gonzalez-Gutiérrez, Mario Castelán</i>	47
Environmental Assessment in Pear Cactus Cultivation Using IoT + edge AI 57 <i>Luis Torres-Treviño, Alejandro Luna-Maldonado, Erik Palacios-Garza, Erick Ordaz-Rivas</i>	57
Traductor de texto audio para ayuda en la enseñanza de débiles visuales 65 <i>Armando Martinez Rios, Adriana Sosa Rojas, Ruben Naftanael Palomino Flores</i>	65
Identificación de patrones lingüísticos, temáticos, emocionales e intención en los corridos y la música regional mexicana 79 <i>José Eder Martínez-Martínez, Yessenia Díaz-Álvarez, Tláloc Humberto Vergara-Morales, Víctor Manuel Millán-Aguilar, Gabriel González-Serna, Noé Alejandro Castro-Sánchez</i>	79
Caracterización de la fatiga mental y su recuperación en tarea de conducción mediante señales electroencefalográficas utilizando algoritmos basados en fractales 93 <i>Diana G. González-Rodríguez, Dulce Martinez-Peon, Erika L. Guzmán Arriaga, Michelle Contró Esparza, Virginia García Pinedo, Carlos A. Rivera Vergara, Gerardo Benavides Bravo, Manuel A. Ramírez Hernández</i>	93

Emulación en FPGA de un memristor y su aplicación en un circuito caótico.....	105
<i>Luis Jothan Gámez Palacios, Opeyemi Micheal Afolabi,</i>	
<i>José Cruz Núñez Pérez</i>	

Early Corruption Detection System Using Machine Learning Algorithms

Maria Fernanda Heredia-Soto, Edgar Huaranga, Franci Suni-Lopez

Universidad de Lima,
Laboratorio de Inteligencia Artificial,
Peru

{mfheredi,ehuarang,fsuni}@ulima.edu.pe

Abstract. This work addresses the development of an early corruption detection system using advanced data analysis and artificial intelligence algorithms. The system uses open data collected from the Anti-Corruption Observatory of Peru’s Comptroller General’s Office. A dataset of 2815 public entities evaluated through the “*Índice de Riesgos de la Corrupción e Inconducta Funcional*” (INCO), with scores ranging from 0 to 100, was structured. The primary task involved classification into six risk levels, from very low to very high. Histogram-based Gradient Boosting (HGB) was applied, achieving an accuracy of 89.6%. Regression tasks on the raw INCO scores and experiments with Large Language Models (LLMs) were also conducted. The system represents an early-stage, yet scalable, tool to support public sector transparency. Future work proposes the integration of explainable AI for improved transparency and real-time policy support.

Keywords: Web scraping, data analysis, histogram-based gradient boosting, machine learning, anti-corruption, early detection model.

1 Introduction

Corruption is a pervasive phenomenon that affects societies at both local and global levels, discrediting public trust, weakening governmental institutions, and undermining sustainable development. It manifests in various forms such as bribery, fraud, extortion, and embezzlement, and is broadly defined by the misuse of official power for personal gain or to benefit third parties. According to the Code of Conduct for Public Officials, corruption involves any act involving illegal benefits in the exercise of public duties. Beyond damaging the integrity of government systems, corruption exacerbates poverty, increases inequality, and violates fundamental human rights [1,2] particularly in less developed democratic systems [3]. International organizations like the United Nations and the Council of Europe stress that national efforts alone are insufficient to tackle this complex issue, emphasizing the need for comprehensive and collaborative strategies to enhance transparency and accountability [4].

In this context, the present study focuses on the creation of an Early Corruption Detection System, leveraging machine learning algorithms to identify high-risk entities proactively. The system is built on open data published annually by the Anti-Corruption Observatory (OBANT) of Peru's Comptroller General's Office (CGR) [12], [13], incorporating artificial intelligence models such as Random Forest and HGB, which are well-suited for detecting patterns in large, heterogeneous datasets [5], [6], [7], [8]. The proposed approach not only aims to detect entities with potential misuse of public resources but also seeks to monitor early warning signals for timely interventions, contributing to improved transparency and efficiency in public management. Furthermore, the initiative aligns with Sustainable Development Goal 16, which promotes strong institutions and the rule of law [9]. The objective of this study is to develop and validate a robust, scalable system using advanced evaluation metrics such as F1-Score, accuracy, and recall to measure its effectiveness.

This paper is organized by outlining the dataset and its relevance, detailing the methodology and results of various machine learning models, discussing key implications for public governance, and concluding with future research directions.

2 INCO Data Description

INCO is a standardized tool created by Peru's Comptroller General's Office to assess the likelihood of corruption and misconduct within public institutions. It is calculated through a composite score derived from multiple officially reported indicators, each assigned a specific weight based on its relevance to corruption risk.

Entities are assigned a score from 0 to 100 based on weighted evaluation of these indicators. For classification, the scores were mapped into six corruption risk levels: Very Low, Low, Moderate, Medium, High, and Very High. The final cleaned dataset included the 2023 INCO scores and the label distribution was as follows: Very High (205), High (248), Medium High (248), Medium (248), Moderate (248), Low (247), and Very Low (100). To address class imbalance, SMOTE (Synthetic Minority Over-sampling Technique) was applied.

The dataset consists of data from 2815 Peruvian state entities, covering the years 2021 to 2023. It includes 24 standardized indicators grouped into two dimensions: "*Institutional Management and Public Integrity*" and "Operational Execution and Control Mechanisms". Key indicators include: (1) sanctions against public officials (I1), (2) administrative/civil/penal responsibilities (I14-I18), (3) irregular public contracting processes (I3, I4, I20), (4) control audit results (I19), and (5) integrity declaration compliance (I13).

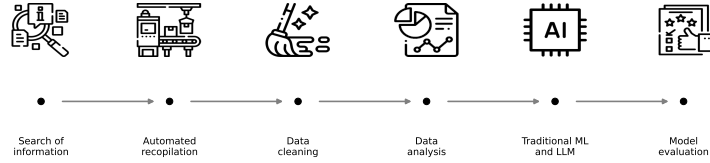


Fig. 1. Workflow of the Methodology.

3 Methodology

This section explains the process for developing the proposed solution. The workflow of the methodology to be followed is shown in Figure 1 from data collection to model evaluation.

3.1 Data Acquisition and Preparation

The primary data source was INCO published by OBANT of CGR [12],[13]. The dataset included information for 2,815 public entities across multiple years (2021–2023), structured around 24 corruption risk indicators. These indicators were grouped into two major dimensions: Institutional Management and Public Integrity and Operational Execution and Control Mechanisms. The features included sanctions against public officials, administrative and penal responsibilities, irregularities in procurement processes, audit results, and compliance with integrity declarations[12],[13].

To build the dataset, relevant information was gathered from public platforms using Web Scraping. Data cleaning involved standardizing variables, removing inconsistencies, and structuring a final dataset ready for machine learning analysis. No manual feature engineering or complex transformations were applied beyond basic consistency adjustments, as the INCO dataset provided pre-aggregated, labeled information suitable for supervised learning.

Feature Analysis To acquire relevant data that contain the necessary information for analysis and detection of potential acts of corruption, sources were identified. Open government data platforms that publish information on public procurement, government expenditures, audit reports, among others, were considered. In some cases, these data were not available on open platforms, so formal requests were also made to obtain the necessary datasets in accordance with transparency laws, allowing their subsequent download. As shown in Figure 2, the most relevant features for predicting corruption risk were selected through an ANOVA-based feature selection process, highlighting indicators related to sanctions, audits, and responsibilities.

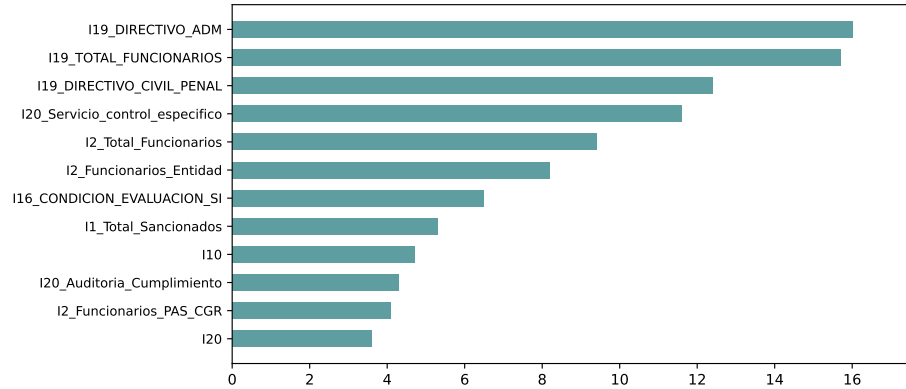


Fig. 2. Features with greater relevance.

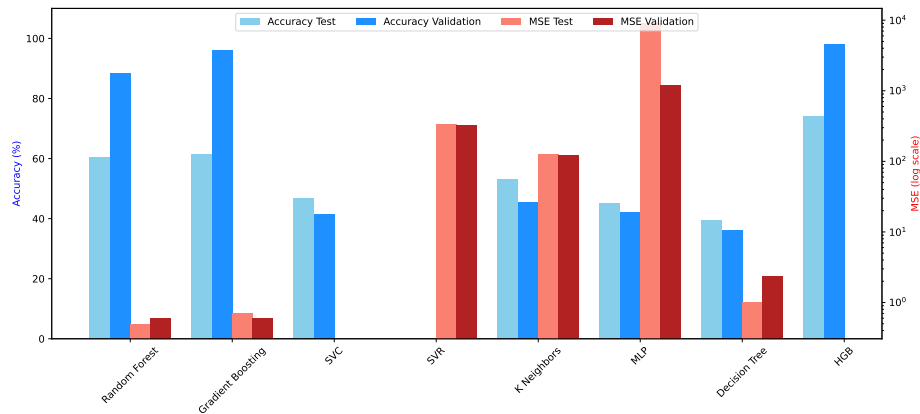


Fig. 3. Modeling.

3.2 Modeling

Two main predictive tasks were addressed: classification and regression. For the classification task, the goal was to categorize entities into six predefined corruption risk levels which are Very Low, Low, Moderate, Medium, High, Very High based on their INCO scores. Several machine learning algorithms were trained and evaluated, including Random Forest, Gradient Boosting, Support Vector Machines (SVM), K-Nearest Neighbors (KNN), and Decision Trees. Among these, the HGB classifier achieved the highest performance and was selected as the primary model. Hyperparameter tuning was conducted using RandomizedSearchCV, and evaluation relied on 5-fold cross-validation.

Figure 3 presents a side-by-side comparison of classification accuracy and regression error (MSE, log-scaled) across multiple models. HGB demonstrates consistently high performance across both tasks, especially in classification

where it significantly outperforms alternatives. Models like MLP and SVC show limitations due to poor handling of class imbalance and feature interactions.

Separately, a regression task was explored to predict the raw INCO scores directly. Models such as Gradient Boosting Regressor and Support Vector Regressor (SVR) were trained; however, regression results did not offer substantial advantages over classification models in terms of interpretability and practical application. In parallel, experiments were conducted using Large Language Models (LLMs) to assess their ability to perform structured corruption risk classification. Models such as Qwen, Llama, Gemma, Hermes, and Mistral were deployed locally via LM Studio. A few-shot prompting strategy was designed, where the LLMs received key indicators as input (e.g., number of sanctioned officials, audit outcomes, pending civil responsibilities) and were asked to predict the corresponding corruption risk level. The outputs were structured to match the classification labels, allowing direct performance comparison with traditional machine learning models.

3.3 Evaluation

Model performance was evaluated using standard classification metrics: Accuracy, F1-Score, and Recall. For regression models, Mean Squared Error (MSE) was used as the primary evaluation criterion. Cross-validation ensured that models generalized effectively and prevented overfitting. The Histogram-Based Gradient Boosting classifier achieved an overall classification accuracy of 89.6%, significantly outperforming the LLM-based approach, which reached 69% accuracy under the same evaluation conditions. The use of Synthetic Minority Over-sampling Technique (SMOTE) contributed to improving detection in minority classes and balancing the overall predictive performance across all corruption risk levels.

4 Results

The features in Figure 2 demonstrated the highest statistical correlation with the target label and were used to guide model training. Their strong relevance underscores the significance of institutional control and sanction-related data in identifying corruption patterns.

The HGB model achieved the best classification accuracy shows in Figure 3 while maintaining low prediction error in regression. Other models, such as SVC and MLP, underperformed due to limited handling of non-linear feature interactions and imbalanced data. These results confirm the robustness of tree-based ensemble methods in heterogeneous public sector datasets.

HGB reached an overall accuracy of 89.6%, with excellent F1-Scores above 98% for the Very High and Very Low classes, where the model's precision and recall were nearly perfect. Table 1 summarizes the classification performance of the HGB and LLM models across corruption risk levels. SMOTE clearly improved detection in minority classes.

Table 1. Classification Results by Risk Level.

Class	Model	Accuracy	Recall	F1-Score	Support
Low	HGB	97.6%	99.2%	98.4%	247
	LLM	74%	74%	74%	100
Moderate	HGB	87.4%	89.1%	88.2%	248
	LLM	64%	64%	64%	150
Medium	HGB	83.1%	81.5%	82.3%	248
	LLM	50%	50%	50%	150
Medium High	HGB	88.9%	87.5%	88.2%	248
	LLM	60%	60%	60%	200
High	HGB	80.7%	81%	80.9%	248
	LLM	65%	65%	65%	205
Very High	HGB	100%	99.6%	99.8%	248
	LLM	80%	80%	80%	100
Accuracy	HGB	-	-	89.6%	1487
	LLM	-	-	69%	905
Macro Avg	HGB	89.6%	89.7%	89.6%	1487
	LLM	68%	68%	68%	905
Weighted Avg	HGB	89.6%	89.6%	89.6%	1487
	LLM	70.5%	70.5%	70.5%	905

Other models like Random Forest and Gradient Boosting (non-HGB) achieved between 88% and 95% accuracy but were more prone to overfitting based on test/validation variance. Support Vector Machines and MLPs underperformed due to data imbalance and lack of deep feature interaction. In the regression task, models achieved moderate results (MSE 0.5–0.7), with no clear added value over classification.

5 Discussion

The empirical results obtained position the Histogram-Based Gradient Boosting (HGB) model as the most effective alternative for the classification of institutional corruption risk based on structured administrative data. Its superior and consistent performance across all predefined risk levels—particularly in the extremes (Very Low and Very High)—underscores its reliability in operational contexts that demand both accuracy and stability. This behavior is indicative of the model’s capacity to capture intricate, non-linear relationships among heterogeneous indicators, such as sanctions, audit outcomes, and declarations of integrity.

The HGB model also demonstrates notable robustness in the presence of class imbalance and missing values, which are common characteristics in real-world public sector datasets. In contrast, algorithms such as Support Vector Classifiers (SVC) and Multilayer Perceptrons (MLP) showed diminished effectiveness, likely

due to their reduced capacity to model inter-variable dependencies and to adapt to uneven label distributions.

Regression-based approaches, while offering continuous predictions, were empirically less effective and substantively less actionable. From an institutional perspective, discrete risk categories are not only more interpretable but also more operationally aligned with audit planning and resource allocation. Furthermore, regression models exhibited higher variance, particularly in mid-level risk predictions, thereby compromising the reliability of their outputs in critical decision-making scenarios.

A conceptual and methodological contrast arises when comparing this work with unsupervised anomaly detection frameworks such as the corruption-based self-supervised model (CAIT)[15]. Whereas CAIT relies on the effect of single-variable corruption on anomaly scores to infer variable relevance, the proposed system utilizes statistical feature selection (ANOVA) and model-specific attribution techniques (e.g., Shapley values) to identify key predictors. This allows for more transparent and reproducible analysis within governance frameworks that require justification and traceability of model outputs.

6 Conclusion

In this paper, we propose a scalable and accurate early corruption detection system for public sector entities, based on machine learning techniques and structured government data. The Histogram-Based Gradient Boosting (HGB) classifier demonstrated the highest performance among all tested models, achieving an overall accuracy of 89.6% and strong F1-Scores across all corruption risk levels. This validates its suitability for institutional monitoring and supports its use in real-world governance applications.

Our results highlight the value of supervised learning algorithms in handling complex, heterogeneous data typical of public administration. While regression models and LLMs were also explored, their practical utility remains limited in this specific classification context. However, future integration of LLMs and explainable AI (XAI) techniques may expand the system's capabilities in terms of interpretability and human oversight.

This approach not only enables proactive identification of high-risk entities but also contributes to policy design, resource allocation, and transparency monitoring. Further work will focus on improving adaptability through real-time data sources, expanding the dataset beyond INCO, and enhancing the model's ability to detect structural patterns of misconduct using causal inference.

Ultimately, the system proposed in this work is a step forward in operationalizing AI for public accountability and offers a practical framework for early detection and intervention in corruption-prone environments.

References

1. Kalienichenko, L. I., Slynko, D. V.: Concept, features and types of corruption. *Law and Safety* 84(1), pp. 39–46 (2022)
2. Oriolo, A.: The Contribution of the European Court of Human Rights to the Construction of a Corruption-Free Society. *International Criminal Law Review* 1, pp. 1–28. Brill Nijhoff (2024)
3. Mubangizi, J.C., Sewpersadh, P.: A human rights-based approach to combating public procurement corruption in Africa. *African Journal of Legal Studies* 10(1), pp. 66–90 (2017)
4. Caruso, S., Bruccoleri, M., Pietrosi, A. : Artificial intelligence to counteract “KPI overload” in business process monitoring: the case of anti-corruption in public organizations. *Business Process Management Journal* {29(4), pp. 1227–1248 (2023)
5. Aguilera-Martos, I., García-Barzana, M., García-Gil, D.: Multi-step histogram based outlier scores for unsupervised anomaly detection: ArcelorMittal engineering dataset case of study. *Neurocomputing* 544, pp. 126228 (2023)
6. Lima, M. S. M., Delen, D.: Predicting and explaining corruption across countries: A machine learning approach. *Government Information Quarterly* 37(1), 101407 (2020)
7. Supriya, M., Adilakshmi, T.: Intrusion Detection in Wireless Sensor Networks Using Histogram Gradient Boosting Classifier. In: *International Conference on Data Science, Machine Learning and Applications*, pp. 473–480. Springer (2023)
8. Rajesh, P.K., Shreyanth, S., Sarveshwaran, R.: Bayesian Optimized Random Forest Classifier for Improved Credit Card Fraud Detection: Overcoming Challenges and Limitations. In: *XVIII International Conference on Data Science and Intelligent Analysis of Information*, pp. 205–214, Springer (2023)
9. Abdelaal, M., Ktitarev, T., Städtler, D. : SAGED: Few-Shot Meta Learning for Tabular Data Error Detection. In: *EDBT*, pp. 386–398 (2024)
10. Köbis, N.C., Starke, C., Edward-Gill, J.: The corruption risks of artificial intelligence. <https://knowledgehub.transparency.org/assets/uploads/kproducts/The-Corruption-Risks-of-Artificial-Intelligence.pdf>, last accessed 2025/03/10
11. Decarolis, F., Giorgiantonio, C.: Corruption red flags in public procurement: new evidence from Italian calls for tenders. *EPJ Data Science* 11(1), 16 (2022)
12. Contraloría General de la República del Perú: Guía Metodológica del Índice de Corrupción e Inconducta Funcional. Observatorio Anticorrupción, Contraloría General de la República del Perú (2023).
13. Contraloría General de la República del Perú: Índice de Riesgos de Corrupción e Inconducta Funcional (INCO) - Guía Metodológica 2024. Contraloría General de la República (2024)
14. Abdelaal, M., Koparde, R., Schoening, H.: Autocure: Automated tabular data curation technique for ML pipelines. In: *Proceedings of the Sixth International Workshop on Exploiting Artificial Intelligence Techniques for Data Management*, pp. 1–11 (2023)
15. Mok, C., Kim, S.B.: Corruption-Based Anomaly Detection and Interpretation in Tabular Data. In: *Pattern Recognition*, 159, 111149 (2024). <https://doi.org/10.1016/j.patcog.2024.111149>

Metodología BRAIN: una metodología para la construcción de chatbots basados en BERT y sus aplicaciones en empresas

Miguel Angel González Reyes, Francisco Edgar Castillo Barrera,
Héctor Gerardo Pérez González

Universidad Autónoma de San Luis Potosí,
Facultad de Ingeniería, San Luis Potosí,
México

a328333@alumnos.uaslp.mx,
{hectorgerardo, ecastillo}@uaslp.mx

Resumen. Este artículo presenta una metodología para la construcción de chatbots basados en el modelo BERT, destacando su aplicabilidad en empresas. Se describe el proceso técnico de desarrollo y se analizan casos de uso en el ámbito empresarial, como la atención de solicitudes de empleo y la provisión de información sobre vacantes. Además, se propone un esquema metodológico que puede ser utilizado como referencia en la implementación de chatbots con inteligencia artificial.

Palabras clave: BERT, NLP, chatbot, inteligencia artificial, empresas, hugging face.

BRAIN Methodology: A methodology for Building BERT-based Chatbots and their Applications in Business

Abstract. This article presents a methodology for building chatbots based on the BERT model, highlighting its applicability in business. The technical development process is described and business use cases are analyzed, such as handling job applications and providing information on vacancies. Furthermore, a methodological framework is proposed that can be used as a reference in the implementation of chatbots with artificial intelligence.

Keywords: BERT, NLP, chatbot, artificial intelligence, business, hugging face.

1. Introducción

En la última década, los avances en inteligencia artificial y procesamiento de lenguaje natural (PLN) han permitido el desarrollo de sistemas conversacionales cada vez más sofisticados. En particular, los modelos basados en arquitecturas transformer

[7], como BERT (Bidirectional Encoder Representations from Transformers), han mostrado resultados sobresalientes en tareas de comprensión y generación de lenguaje.

Este artículo presenta BRAIN (Base Extraction, Refinement, Answering, Information Retrieval, Natural Interaction), una metodología paso a paso para la construcción de chatbots basada en BERT, diseñada específicamente para aplicaciones en entornos empresariales. La metodología busca integrar herramientas como Hugging Face [6], SpaCy [8] y Sentence Transformers [7] para permitir una implementación eficiente, precisa y adaptable de chatbots que puedan interactuar con documentos institucionales, como vacantes o manuales internos.

A lo largo del trabajo se detallan las cinco fases de la metodología BRAIN, su implementación en Python, los resultados obtenidos en un caso de uso real, y un análisis comparativo respecto a metodologías existentes. Además, se presentan métricas de evaluación y una discusión sobre el impacto y las oportunidades de mejora del enfoque propuesto.

2. Metodología BRAIN

La metodología propuesta se compone de cinco fases secuenciales que permiten transformar información documental en respuestas conversacionales contextualizadas. A continuación, se describe en detalle cada componente de la metodología BRAIN:

2.1 Base data extraction

Esta fase se enfoca en la recopilación y extracción del conocimiento fuente desde archivos institucionales. En el caso de uso, se procesaron documentos PDF con descripciones de puestos y requisitos laborales mediante la biblioteca PyPDF2. El texto se segmentó por párrafos y se descartaron líneas vacías o irrelevantes. Esta información constituye el corpus base sobre el cual se construye el modelo conversacional.

2.2 Refinement of text

Una vez extraída la información, se procede a su limpieza y normalización. Se utilizaron técnicas de normalización como la eliminación de acentos y conversión a minúsculas. Luego, se aplicó el modelo preentrenado all-MiniLM-L6-v2 de Sentence Transformers [7] para convertir cada párrafo en un vector semántico. Esto permite evaluar similitudes conceptuales más allá de la coincidencia literal de palabras.

2.3 Answer model with BERT

Se utilizó el modelo de Hugging Face [6], para implementar una arquitectura de pregunta-respuesta en español. Aunque en la versión actual del sistema el modelo no se invoca directamente para generar respuestas, su incorporación está planeada como una mejora futura para afinar la selección de fragmentos contextuales.

2.4 Information retrieval strategy

El sistema implementa una estrategia híbrida para recuperar la información más relevante:

Coincidencia exacta. Si una pregunta coincide textualmente con algún fragmento del corpus, se recupera esa sección.

Heurística por intersección de palabras clave. si no hay coincidencia exacta, se extraen palabras clave de la pregunta y se comparan con cada párrafo del corpus. El párrafo con mayor número de palabras coincidentes se selecciona como candidato.

Esta técnica permite responder incluso a preguntas mal formuladas o incompletas, maximizando la utilidad del chatbot en situaciones reales.

2.5 Natural interaction implementation

Finalmente, se desarrolló una interfaz interactiva en Python donde los usuarios pueden formular preguntas en lenguaje natural. La interacción se realiza mediante consola, permitiendo que el sistema procese la entrada en tiempo real y devuelva una respuesta basada en los criterios anteriores. Se prevé extender esta funcionalidad mediante interfaces web o APIs REST para mejorar su usabilidad en entornos empresariales.

3. Resultados y evaluación

Aunque la metodología BRAIN fue concebida originalmente para su uso en contextos empresariales (por ejemplo, atención a candidatos en procesos de reclutamiento), su primera implementación funcional se utilizó en un dominio académico. El objetivo fue validar su rendimiento empleando como dominio temático la orientación vocacional para aspirantes a los programas educativos que se ofrecen en la Universidad Autónoma de San Luis Potosí (UASLP).

3.1 Contexto de la evaluación

Se utilizó como fuente de conocimiento un documento PDF fabricado a partir de documentación oficial que describe dos programas educativos: Ingeniería en Computación e Ingeniería en Sistemas Inteligentes. El texto incluye preguntas frecuentes sobre áreas de énfasis, perfil de ingreso y egreso, oportunidades laborales, movilidad y requisitos de admisión.

3.2 Metodología de evaluación

Se realizaron 15 pruebas consecutivas, simulando preguntas frecuentes que podrían formular aspirantes reales sin conocimientos técnicos, incorporando intencionalmente lenguaje coloquial característico de una conversación casual y errores ortográficos, con el fin de verificar el funcionamiento del sistema.



Fig. 1. Metodología BRAIN, representada como una secuencia de pasos para la creación de chatbots basados en el modelo BERT.

3.3 Criterio de respuesta

El sistema procesa la entrada buscando:

- Coincidencia exacta del texto en los párrafos normalizados.
- Si no encuentra coincidencia, emplea una estrategia de intersección de palabras clave.
- En ambos casos, devuelve el primer párrafo completo hasta el punto final (.) en el texto original.

3.4 Ejemplos de interacción simulada

3.5 Observaciones

- El sistema fue capaz de entregar respuestas relevantes en 14 de las 15 pruebas realizadas.
- Las preguntas formuladas con errores ortográficos o sin signos fueron correctamente entendidas.
- Las respuestas son literales del texto fuente, lo que garantiza su precisión.
- El sistema es especialmente eficaz para consultas concretas, pero podría mejorarse en respuestas abiertas implementando un sistema de validación por refuerzo.

4. Discusión

La implementación de la metodología BRAIN en un caso académico permite reflexionar sobre su eficacia, versatilidad y posibles limitaciones. El objetivo original del proyecto fue diseñar un chatbot de asistencia académica, orientado a apoyar a aspirantes universitarios en su proceso de orientación vocacional. Por ello, las pruebas

Table 1. Ejemplos de búsquedas realizadas con alumnos de nuevo ingreso que les hubiera gustado poder consultar sus dudas cuando eran aspirantes universitarios en la UASLP

Pregunta del usuario.	Respuesta obtenida
que ase un ingeniero en sistemas	Un ingeniero aplica conocimientos científicos y tecnológicos para resolver problemas.
que es un sistema inteligente	Los sistemas inteligentes son aquellos que perciben y responden al mundo que los rodea utilizando dispositivos computacionales.
hay trabajo para esa carrera	¡Sí, bastante! Hay una demanda muy grande a nivel mundial de profesionales de la computación, y con la llegada de la Industria 4.0, se requieren cada vez más ingenieros capaces de crear sistemas inteligentes.
en donde se puede trabajar	¡En donde tú quieras! A diferencia de otras carreras, las computadoras están y se necesitan en todos lados.
que materias tiene esa carrera	La disciplina de la computación es cada vez más amplia y diversa, con un rango de retos y oportunidades, por lo que hemos definido cuatro áreas de énfasis (o “tracks”), en donde te puedes especializar en ciertos temas en los últimos dos años de tu carrera.
necesito saber programar antes de entrar	No es necesario que sepas programar, pero si nunca lo has hecho te recomendamos que investigues lo que es la programación e intentes realizar algún programa sencillo, para que te des una idea de lo que implica.
como puedo estudiar en esa carrera	Para entrar a esta carrera deberás cumplir con los siguientes requisitos: Bachillerato completo en Ciencias Físico-Matemáticas, Bachillerato Tecnológico, Bachillerato General o Único.
que perfil necesito para entrar	Para ingresar y tener éxito en esta carrera es deseable que cuentes con: Razonamiento lógico y matemático Capacidad de iniciativa y liderazgo Capacidad de abstracción y síntesis Capacidad de razonamiento, análisis y comprensión de lectura.
cuales son los puestos donde puedo trabajar	Como ingeniero en sistemas inteligentes serás capaz de diseñar y desarrollar nuevos sistemas computacionales, por lo tanto, algunos puestos en donde te podrías desempeñar son: programador, analista de sistemas, administrador de redes.
que diferencia hay con otra ingeniería	La diferencia entre estas dos carreras es que tienen materias diferentes y enfoques diferentes, sistemas se basa en programación y computación en hardware.
que necesito para llegar a ser un ingeniero experimentado	Debes de practicar mucho.

de validación se enfocaron exclusivamente en este dominio, utilizando documentación oficial de programas educativos como fuente de conocimiento.

Durante el desarrollo de BRAIN, se identificó que la estructura modular propuesta —basada en la extracción de información, normalización semántica y recuperación híbrida de respuestas— tiene el potencial de adaptarse también a otros contextos más allá del académico. En particular, se reconoció su aplicabilidad en entornos empresariales para tareas como *la atención de solicitudes de empleo, programas de capacitación interna o soporte al cliente*, lo cual abre oportunidades de expansión futura para la metodología.

Uno de los principales hallazgos de la evaluación fue que la estrategia híbrida de recuperación de información —que combina coincidencia exacta e intersección de palabras clave— resultó altamente eficaz para responder preguntas comunes, incluso cuando estas contenían errores ortográficos o formulaciones ambiguas. Esto sugiere

Table 2. Tabla comparativa sobre proyectos recientes relacionados a chatbots con BERT y sus diferencias con la metodología BRAIN.

Referencia	¿Usa Metodología Formal?	Técnica Principal	Diferencias frente a BRAIN
Gopisetty (2024) [1]	No	Chatbot basado en BERT para e-commerce	Implementa sistema pregunta-respuesta para productos, pero no estructura modular definida; requiere entrenamiento específico para cada dominio. BRAIN es más generalizable y modular.
Babu & Boddu (2024) [2]	Parcial	Chatbot médico usando BERT fine-tuned	Utiliza fine-tuning en dominio médico, lo cual requiere datasets anotados. BRAIN evita fine-tuning pesado usando embeddings semánticos y búsqueda híbrida.
Singh et al. (2024) [3]	No	BERT combinado con Aprendizaje por Refuerzo (RL)	Enfocado en optimizar interacciones vía RL, no describe una metodología paso a paso; mayor complejidad técnica. BRAIN es más sencillo y directo para implementación práctica.
Hafidz et al. (2024) [4]	No	Chatbot de información turística con BERT	Utiliza BERT como motor de respuestas, pero sin un flujo estructurado ni estrategias híbridas de recuperación de información. BRAIN define un proceso modular claro.
Zeniarja et al. (2025) [5]	Parcial	Chatbot de salud mental usando BERT + ELECTRA	Integra dos modelos para mejorar comprensión emocional; enfoque de alta complejidad técnica. BRAIN está orientado a facilitar despliegues rápidos en contextos documentales bien definidos.

que, en dominios bien estructurados, sistemas basados en búsqueda semántica pueden ofrecer resultados satisfactorios sin necesidad de recurrir a modelos generativos complejos.

El uso de herramientas accesibles como Hugging Face [6], SpaCy [8] y Sentence Transformers [7] facilitó el desarrollo del sistema sin requerir grandes volúmenes de datos ni entrenamientos específicos, lo que disminuye las barreras de entrada para instituciones educativas o pequeñas organizaciones que buscan soluciones de atención automatizada.

No obstante, se reconocen algunas áreas de mejora. El sistema muestra limitaciones en la respuesta a preguntas abiertas o de carácter subjetivo, debido a su dependencia de fragmentos literales extraídos del corpus fuente. Futuras extensiones del sistema podrían considerar la incorporación de modelos BERT en la fase de refinamiento de respuestas, así como el uso de técnicas de aprendizaje por refuerzo para incrementar la precisión y adaptabilidad a distintas intenciones del usuario.

Finalmente, aunque los resultados preliminares en el dominio académico fueron positivos, será necesario realizar evaluaciones más amplias, incluyendo pruebas con usuarios con diferentes intereses y en diferentes contextos, para validar la efectividad y escalabilidad de la metodología BRAIN en entornos diversos.

Para reforzar la originalidad y aporte de la metodología BRAIN, se presenta una comparación con trabajos recientes que implementan chatbots basados en modelos BERT, pero que no necesariamente estructuran una metodología modular clara. La Tabla 2 muestra las principales diferencias:

Como puede observarse en la tabla 2, BRAIN muestra ventajas por su diseño modular, su accesibilidad para implementación sin entrenamiento intensivo y su estrategia híbrida de recuperación de información, lo que la hace especialmente

adecuada para contextos donde se dispone de corpus documentales definidos, pero no de grandes datasets anotados.

5. Conclusiones

La metodología BRAIN representa una propuesta estructurada, accesible y efectiva para el desarrollo de chatbots basados en modelos de lenguaje preentrenados como BERT. Su enfoque modular facilita la comprensión, implementación y adaptación en distintos dominios, ya sea empresarial, académico o institucional.

A través del caso de estudio presentado —un chatbot orientado a resolver dudas frecuentes de aspirantes universitarios sobre carreras en computación— se demostró que BRAIN permite construir sistemas funcionales, capaces de entregar respuestas relevantes sin necesidad de grandes volúmenes de datos ni entrenamiento supervisado. El sistema logró una tasa de éxito alta al recuperar información útil incluso en condiciones de formulación imperfecta por parte del usuario.

Se concluye que:

- El uso de modelos de embeddings [9] y estrategias de recuperación híbridas [10] puede ser suficiente para tareas de respuesta automática cuando el contexto está bien definido.
- La arquitectura propuesta permite extender la funcionalidad fácilmente, ya sea con nuevas fuentes de conocimiento o con modelos más sofisticados para el refinamiento de respuestas.
- Su diseño la convierte en una herramienta formativa útil, tanto para enseñar fundamentos de procesamiento de lenguaje natural como para desarrollar prototipos de soluciones conversacionales en el mundo real.

Como trabajo futuro, se buscará añadir al sistema que maneje BERT directamente, junto al sistema híbrido que ya se tiene para buscar obtener mejores resultados, se recomienda continuar con la validación de la metodología en escenarios más diversos, incluyendo interacciones con usuarios interesados en el contexto en que se quiera usar el ChatBot, así como explorar su integración con interfaces web y APIs para ampliar su aplicabilidad en entornos de producción.

Agradecimientos. Agradezco al Profesor Dr. Juan Carlos Cuevas Tello por sus enseñanzas y por sentar las bases de conocimiento que fueron fundamentales para el desarrollo de esta investigación.

Referencias

1. Gopisetty, G., *Et al.*: Chatbot Building with BERT for E-Commerce. International Journal for Research in Applied Science and Engineering Technology, (2024) doi: 10.22214/ijraset.2024.59449.

2. Babu, A., Boddu, S.: BERT-Based Medical Chatbot: Enhancing Healthcare Communication Through Natural Language Understanding. *Exploratory Research in Clinical and Social Pharmacy*, 13. (2024) doi: 10.1016/j.rcsop.2024.100419.
3. Praneeth, KR., Ruprah, T.S., Madhuri, J.N.: Optimizing Customer Interactions: A BERT and Reinforcement Learning Hybrid Approach to Chatbot Development. *International Journal of Advanced Computer Science and Applications*, 15(9), (2024)
4. Hafidz, I., Mukti, B., Zahra, Q.: Chatbot Model Development Using BERT for West Sumatera Halal Tourism Information. *Halal Research Journal*, 4(2), (2024) doi: 10.12962/j22759970.v4i2.1819.
5. Zeniarja, J., Paramita, C., Subhiyakto, E.: Integrating ELECTRA and BERT Models in Transformer-Based Mental Healthcare Chatbot. *Indonesian Journal of Electrical Engineering and Computer Science*, 37(1), pp. 315–314 (2025) doi: 10.11591/ijeecs.v37.i1.pp315–324.
6. Osborne, C., Ding, J., Kirk, H.R.: The AI Community Building the Future? A Quantitative Analysis of Development Activity on Hugging Face Hub. *J Comput Soc Sc* 7, pp. 2067–2105 (2024) doi: 10.1007/s42001-024-00300-8.
7. Reimers, N., Gurevych, I.: Sentence-BERT: Sentence Embeddings Using Siamese BERT-Networks. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing*, pp. 3982–3992 (2019) doi: 10.18653/v1/D19-1410.
8. Honnibal, M., Montani, I., Van Landeghem, S.: spaCy: Industrial-Strength Natural Language Processing in Python, (2020) doi: 10.5281/zenodo.1212303.
9. Mikolov, T., Chen, K., Corrado, G.: Efficient Estimation of Word Representations in Vector Space. *Proceeding of Workshop at ICLR*, (2013) doi: 10.48550/arXiv.1301.3781.
10. Karpukhin, V., Oguz, B., Min, S.: Dense Passage Retrieval for Open-Domain Question Answering. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing*, pp. 6769–6781 (2020) doi: 10.18653/v1/2020.emnlp-main.550.

Transformer Architecture Reconfiguration Using Heterogeneous Attention

Xamanek Martínez-Marín, José Adan Hernández-Nolasco,
Noel Zacarias Morales

Universidad Juárez Autónoma de Tabasco,
División Académica de Ciencias y Tecnologías de la Información,
México

xamanekmtz@gmail.com, adan.hernandez@ujat.mx,
n1.zkmt@gmail.com

Abstract. Each attention mechanism in neural network architectures specializes in capturing specific and relevant aspects of input data sequences, allowing the model to focus on the most significant parts for the task at hand. However, this focus is performed in a "homogeneous" manner; that is, the model applies a uniform approach across all training epochs, which can result in output tensors containing incomplete or biased information. This research explores a new direction by proposing the incorporation of various attention mechanisms within a single model to create what we call a "heterogeneous" attention mechanism. For this purpose, heterogeneous attention mechanisms (HAM) A and B are implemented, which make use of "Soft Attention" as a complementary mechanism and integrate the output tensors using "Hadamard Product" and "Element-Wise Addition" respectively. The Vision Transformer (ViT) is used in image classification as a task to evaluate the performance of three models, ViT HAM A, ViT HAM B and ViT Classic. The results show an improvement in convergence time of almost 9.5% and better accuracy of ViT HAM B with respect to ViT Classic. This research opens the door to new methods to improve the efficiency and effectiveness in the training of deep learning models through the use of diversified attention mechanisms.

Keywords: Attention mechanism, neural networks, deep learning, transformer.

1 Introduction

The introduction and continuous improvement of attention mechanisms has led to significant advances in the field of artificial intelligence, laying the groundwork for the development of more advanced architectures, such as Transformers, which rely exclusively on attention mechanisms to efficiently and effectively process data streams. These advances have revolutionized not only machine translation but also a wide range of applications in natural language processing (NLP),

computer vision, and beyond, making the attention mechanism one of the most influential innovations in the recent history of machine learning (ML).

Given the computational complexity of the Transformer, the goal is to make efficient use of the resources used for training, so reducing the total training time with minimal hardware is an attractive option. Taking this into account, an architectural improvement process is carried out, proposing a variant of the attention mechanism applied in deep learning research.

The traditional Transformer architecture employs an attention mechanism called "Self Attention," which performs "homogeneous" focusing, that is, it applies a uniform focus across all training epochs, which can result in output tensors containing possibly incomplete or biased information. It is proposed to incorporate various attention mechanisms within a single model to create what we call a "heterogeneous" attention mechanism. The central idea is that, by integrating multiple types of attention that operate in a complementary manner, the model can capture relevant information in a more diverse and complete way during each training epoch. This variability in the focus of attention allows the model to avoid excessive dependence on certain features, thus improving the quality of the representation of the input data. Furthermore, it is proposed that this heterogeneous approach not only improves the model's ability to generalize, but also reduces the overall training time. By capturing relevant information more efficiently in each iteration, the model requires fewer epochs to converge, optimizing the use of computational resources and making the training process more efficient in terms of time and processing.

This study focuses on analyzing the performance of the proposed attention mechanism using the Vision Transformer architecture in image classification tasks using the CIFAR-10 dataset. Evaluation in other visual domains and a comprehensive comparison with traditional convolutional architectures are not considered. The scope is limited to analyzing the model's accuracy, computational efficiency, and ability to visually identify the regions representing the subject of interest in the images within the defined experimental context. Future research could extend this approach to segmentation or object detection tasks in more complex datasets and in other domains such as natural language processing (NLP) tasks.

2 Related Works

The attention mechanism has evolved and diversified into various implementations that have demonstrated different degrees of complexity, efficiency and effectiveness, in 2017 Vaswani et al [1] implemented the Transformer model that makes use of the "Self Attention" attention mechanism to calculate the attention scores of the sequence, it has also been implemented in the Vision Transformer (ViT) in 2020 by Alexey Dosovitskiy et al [2] to do image recognition.

This made "Self Attention" the most famous attention mechanism today, Vaswani introduced it in his work and replaces recurrent networks with this

mechanism as the basis for sequential data processing; it has been the basis of current natural language processing models such as GPT, BERT among others. Various works have been derived from the transformer that have sought to make Vaswani's Transformer more efficient, for example the Linformer [3] that reduces the complexity to $O(n)$ by approximating the attention matrices, the Longformer [4] being very similar to the standard Transformer introduces a combination of local and global attention to reduce the complexity to $O(n)$ in long sequences, the Performer [5] makes use of an approximation technique based on kernel functions called FAVOR+ (Fast Attention Via positive Orthogonal Random features) with which it reduces the complexity of the attention calculation to $O(n)$.

3 Methodology

3.1 Heterogeneous Attention Mechanisms

The proposed mechanism is based on the premise that multiple viewpoints can provide a more robust and effective solution to a problem than a single approach could offer. In the context of attention mechanisms within neural networks, this idea translates into the integration of different types of attention, such as "soft attention" [6], "global attention" [7], "local attention" [7], and so on, which act as complementary perspectives. These perspectives support the "self-attention" mechanism, allowing the model to not only focus on global relationships within the input sequence, but also capture finer details and contextual nuances that could be essential for a more complete and accurate understanding of the data.

Vision Transformer. A ViT is implemented with the following characteristics in its architecture and as shown in Fig 1 (Classic ViT Diagram).

3.2 Heterogeneous Attention Mechanism A (HAM A)

This mechanism implements "soft attention" (See Fig 2) as a complementary mechanism to "self attention" and integrates them into one using Hadamard Product [8], which consists of multiplying each element of the first tensor with the corresponding element of the second tensor at the same position, this operation is usually denoted as:

$$C = A \odot B,$$

where:

$$\begin{aligned} A, B &\in \mathbb{R}^{m \times n}, \\ C_{ij} &= A_{ij} \cdot B_{ij}. \end{aligned}$$

Hadamard product does not mix cross-information between different positions, it only locally affects each element.

The incorporation of the attention mechanism is done within the attention head which is within the Multi-Head Head (MHA) which is found within each of the four blocks that make up the Encoder and at the end of which the "soft attention" output tensor is integrated with the "self attention" output tensor.

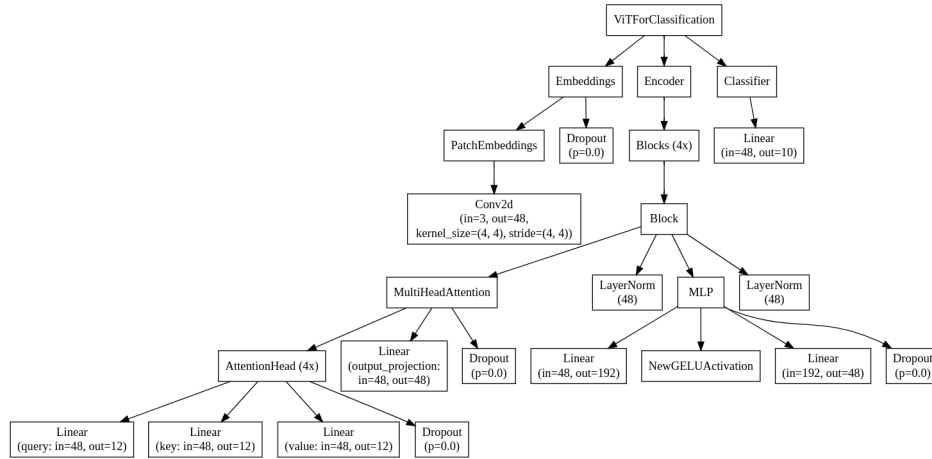


Fig. 1. Structural diagram of the Classic Vision Transformer (ViT) model used for image classification tasks.

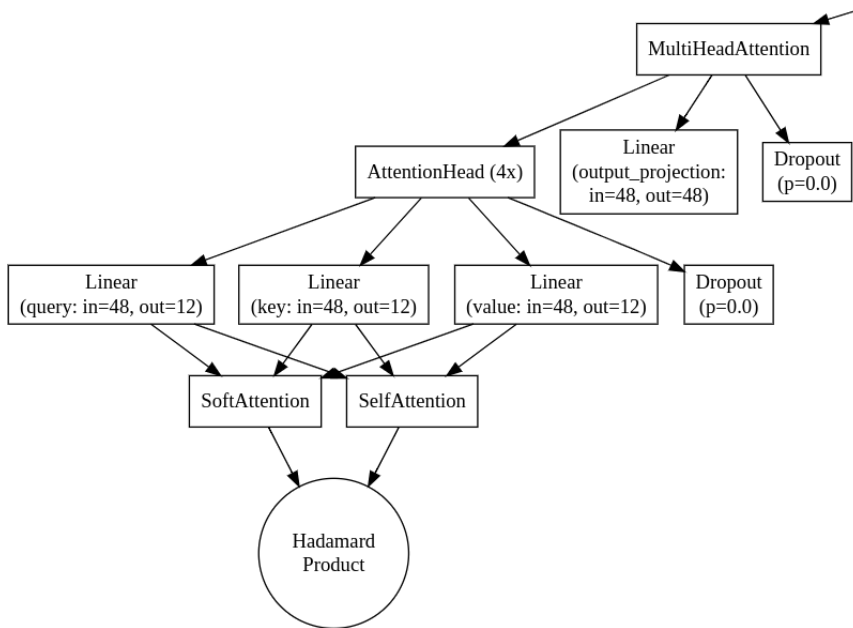


Fig. 2. Structural diagram of the Vision Transformer HAM A (ViT) model used for image classification tasks.

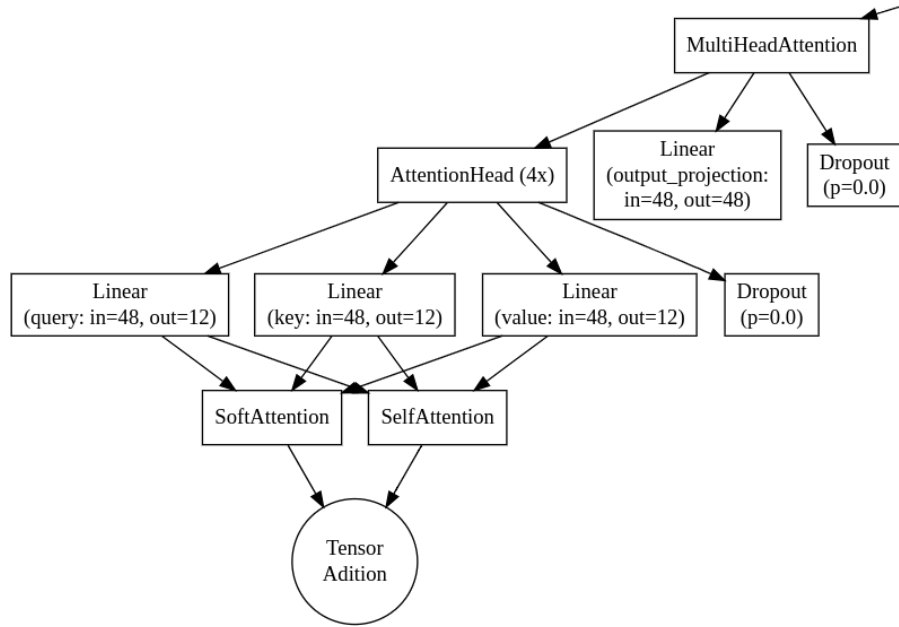


Fig. 3. Structural diagram of the Vision Transformer HAM B (ViT) model used for image classification tasks.

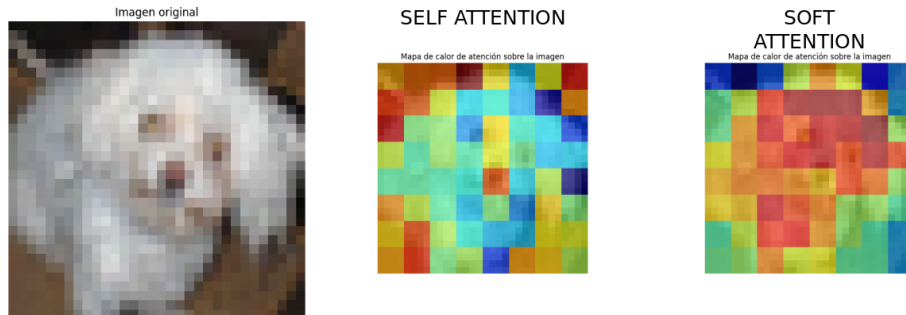


Fig. 4. Comparison between the original image (CIFAR-10) and attention maps of each mechanism separately.

3.3 Heterogeneous Attention Mechanism B

This mechanism implements "soft attention" (See Fig 3) as a complementary mechanism, but integrates them with a one-to-one addition of tensors (Element-Wise Addition) [9], which is a mathematical operation in which two tensors of the same size are added position by position, each element of the resulting tensor is the sum of the corresponding elements of the original tensors.

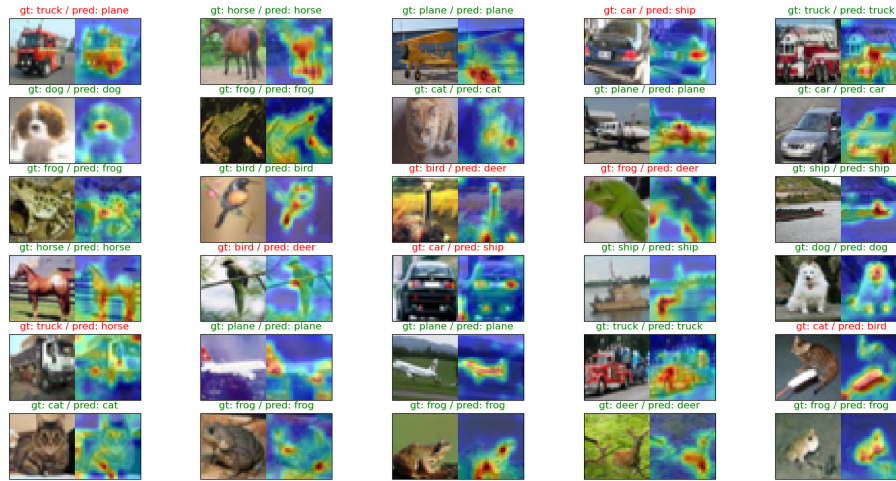


Fig. 5. Heat map generated by the attention captured by the Classic ViT model after training.

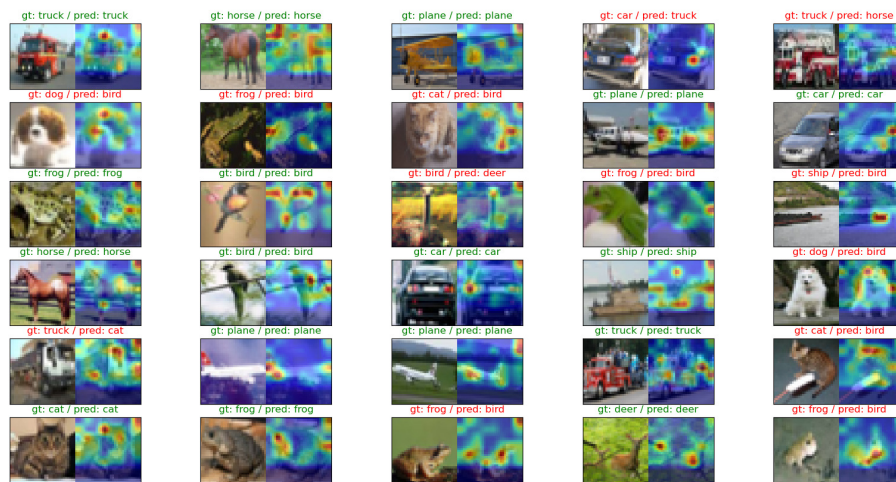


Fig. 6. Heat map generated by the ViT HAM A model’s captured attention after training.

Given two tensors A and B of the same shape $m \times n$ their element-by-element sum is defined as:

$$C_{ij} = A_{ij} + B_{ij},$$



Fig. 7. Heat map generated by the ViT HAM B model's captured attention after training.

the result C is also a tensor of shape $m \times n$. As can be seen in Fig 3, the configuration of the ViT architecture is the same as the classical model and model A, with the difference being the integration of the output tensors of both mechanisms.

3.4 Pytorch

For the implementation and training of ViT, Pytorch (v2.3.0) is used, which is an open source framework for numerical computing and machine learning developed by Meta AI Research [10], Pytorch allows defining DL models such as ViT from scratch or using predefined models, it is possible to modify the number of classes, size of the patches, number of layers, etc. It allows the preprocessing of images by transforming them into the appropriate format, resizing, conversion to tensors and normalization.

Using built-in tools like "DataLoader" and "Dataset" allows you to load images in mini-batches and apply automatic transformations during training. It also enables the ability to define a loss function and an optimizer. It also allows us to use specialized hardware (GPUs) and local distributed computing (multiple GPUs or multiple computers on a network).

For this case, a ViT was implemented from scratch, Figs 1, 2 and 3 were created from the output of the "print(model)" command.

3.5 Dataset

For the image classification task in which the training and subsequent evaluation of the trained models was carried out, the CIFAR-10 Dataset is used, which was developed by Alex Krizhsky, Vinod Nair and Geoffrey Hinton [11] at the University of Toronto as part of the deep learning in computer vision project. This dataset has a total of 60,000 images that are grouped into 10 classes with a total of 6,000 images per class, at a resolution of 32 x 32 pixels and 3 channels (RGB), the storage format used was the binary matrix.

The dataset is divided into 50,000 images for training and 10,000 for testing, each set of images from all classes with a uniform distribution. The 10 object classes are "airplane", "automobile", "bird", "cat", "deer", "dog", "frog", "horse", "ship", "truck", the dataset presents an intermediate difficulty given the small image size with a lot of background noise and visual variability, this represents a reasonable challenge for simple models and a good benchmark for modern architectures such as ViT.

3.6 Hyperparameters

These are external configurations to the model that are not directly learned during training, but control the behavior of the learning process. Unlike model parameters (such as layer weights), hyperparameters must be defined before training and affect both the architecture and the dynamics of the learning process.

For ViT training, an optimized architecture was configured to operate with low-resolution images, specifically the CIFAR-10 Dataset, whose dimensions are 32 x 32 pixels and contain three channels corresponding to the RGB format. The implemented architecture is characterized by a patch size of 4 x 4, which allows each image to be represented by a sequence of 64 tokens. The size of the hidden embedding was set to 48 dimensions, and the model was structured with four layers of multi-head attention blocks, each with four parallel attention heads. Inside each block, a multi-layer perceptron was used with an intermediate dimension of 192 nodes, which corresponds to four times the size of the embedding. The GELU activation function was used, commonly used in Transformer-type models due to its smoothness and nonlinear modeling capacity (See Table 1).

Regarding image preprocessing, standard computer vision transformations were applied to enhance the model's generalization capabilities. The images were resized and subsequently subjected to random crops varying in both scale and aspect ratio, followed by random horizontal mirroring with a 50% probability. These techniques, which are aligned with data augmentation strategies, seek to simulate natural variations in the data without altering its semantics. Finally, the tensors were normalized using the mean and standard deviation of the CIFAR-10 set, which contributes to stabilizing the learning process by centering the distribution of the input values (See Table 2).

This set of hyperparameters and transformations was selected to balance the model’s representational capability with computational efficiency, making training feasible in resource-limited environments without sacrificing the model’s ability to capture relevant patterns in the data.

The model was trained for 400 epochs using the AdamW optimizer, known for its effectiveness in combining weight decay techniques with adaptive update steps. The learning rate was set to 0.01, a value that, although high in absolute terms, can be adequately managed using learning programming strategies such as progressive decay or warm-up schemes, although neither was used in this case. A batch size of 256 samples per iteration was used, which favors a more stable estimation of the gradients, especially in the presence of large-batch regularization techniques.

4 Results

In Fig 4 we can observe a visual comparison between the original input image randomly extracted from the CIFAR-10 dataset and the attention maps generated by ViT in a single step using a separate attention mechanism.

Table 1. Hyperparameters used for training the Vision Transformer model.

Hyperparameter	Value (applied style)
Image size	32×32
Input channels	3
Number of classes	10
Patch Size	4
Embedding Size	48
Layers	4
Attention Heads	4
Intermediate MLP size	192
Attention Dropout	0.0
Hidden Dropout	0.0
Initialization of weights	0.02
Activation function	GELU
Learning rate	0.01
Batch Size	256
Epochs	400
Optimizer	AdamW

In the Self Attention map, each token (image patch) directly attends to every other token using learned attention weights, allowing for global interactions. Regions highlighted in red are primarily boundaries of the dog’s body or very sharp contrasts (such as a dark background versus a white coat); these contrasts can usually be detected as useful discriminants and therefore receive high attention weights. Unlike convolutional neural networks (CNNs) [12], ViTs do

Table 2. Transformations applied during image preprocessing.

Transformation	Parameters (applied style)
Random Crop	Tamaño: 32, Padding: 4
Random Horizontal Flip	Probability: 0.5
Random Resized Crop	Scale: (0.8, 1.0), Ratio: (0.75, 1.333)
Resize	32 × 32
ToTensor	—
Normalize	Mean: (0.4914, 0.4822, 0.4465), Standard Deviation: (0.2023, 0.1994, 0.2010)

Table 3. Conceptual comparison between attention mechanisms in Vision Transformer.

Criterion	Self-Attention	Soft Attention
Spatial dispersion	<i>High, with peaks at the edges</i>	<i>Low, concentrated in the center</i>
Multiple foci of attention	<i>Yes, can serve multiple regions</i>	<i>No, primarily unimodal focus</i>
Contextual Flexibility	<i>Global, allows full interaction between tokens</i>	<i>Partial, with focus on a dominant region</i>
Interpretability	<i>Complex but rich in spatial relationships</i>	<i>Clear, with easily interpretable localized focus</i>
Observed pattern	<i>Dispersed attention including periphery</i>	<i>Focused circular attention with smooth decay</i>

Table 4. Comparative results of accuracy and loss in training and testing for three variants of the ViT model.

Name	Best Accuracy & epoch	Last Accuracy	Best Train Loss & epoch	Last Train Loss	Best Test Loss & epoch	algo borre
ViT Classic	0.7977	0.7902	0.4863	0.5011	0.6123	0.6156
	Epoch 377	Epoch 400	Epoch 390	Epoch 400	Epoch 377	Epoch 400
ViT HAM A	0.7864	0.7723	0.5275	0.5424	0.6165	0.6641
	Epoch 393	Epoch 400	Epoch 398	Epoch 400	Epoch 393	Epoch 400
ViT HAM B	0.8107	0.7693	0.5015	0.5130	0.5558	0.6882
	Epoch 342	Epoch 400	Epoch 396	Epoch 400	Epoch 342	Epoch 400

not have an inductive bias towards local structures, meaning that the network does not automatically favor the center or assume that nearby regions are related. Therefore, it is perfectly plausible that distant areas can receive high attention if they help build a useful representation.

It is emphasized that this attention map corresponds to a single step without additional processing; these early layers tend to learn low-level patterns or global geometric structure. The attention to edges reflects that at this stage, the network has not yet refined its focus toward the central semantic region (the dog's face). This is a natural behavior of the "Self Attention" mechanism.

The "Soft Attention" attention map (right in Fig 4) weighs different regions of the same image, with attention strongly concentrated in the center of the image, particularly in the areas of the dog's face, which includes the snout, eyes, and part of the forehead. Unlike the "Self Attention" attention map, here we observe a gradual and symmetrical decrease in attention towards the edges of

the image; the corners and margins present cooler colors, indicating that the model does not consider them "informative" in this first and only step. The "circular" shape of the attention is reminiscent of a "Gaussian window" type approach centered on the region of greatest visual density (representative of the subject in the image).

This behavior is due to the design of the attention mechanism, which generally consists of a normalized distribution of weights (via Softmax) applied to spatial or embedded features, which forces attention to focus strongly on a few regions and quickly decay outside of them.

This map reflects how "Soft Attention" prioritizes the central region of the image (where relevant semantics are located) in a more direct and localized way than "Self Attention." This central attention will remain constant across multiple layers, reinforcing classification based on a salient area (and reinforcing the attention captured by "Self Attention").

Table 3 shows a conceptual comparison between both attention mechanisms according to the previous analysis.

In Table 4 we have the ViT HAM B implementation reaching a maximum accuracy of 0.8107 at epoch 342 being higher than the other two ViT implementations, indicating that HAM B can capture patterns more effectively at some point during its training, this makes it the most accurate model during this stage making it a more efficient implementation in terms of convergence, reaching its peak performance before Classic and HAM A, however we see a considerable drop in the last epoch indicating overfitting at the end of training.

The attention shown by the heat maps in figures 5, 6 and 7 indicate that the "HAM B" model focuses more globally to capture image features, suggesting that it has a better ability to capture specific details or patterns in some examples, reflected in its better accuracy in some categories. In the images we observe a better focusing on the correct areas (for example, in the "frog" and "bird" images) this is a sign that the implemented attention mechanism has a better performance in terms of prioritizing the most relevant features in those classes.

5 Conclusions

The poor performance of the HAM A implementation compared to the other two models is mainly due to the fact that tensor integration using the Hadamard Product is not reversible, meaning that different combinations of values can give the same product (e.g., $2 \odot 3 = 6$, as well as $1 \odot 6$). Information is lost if one of the elements is zero, which causes the contribution of the other tensor to be completely lost, and finally, it does not preserve the original scales of the tensors.

The overfitting present in the HAM B model in its last epoch can be mitigated using several techniques in this type of tasks, starting by adding more "data augmentation", other transformations that were not initially included such as Vertical Flip, CutMix or MixUp can be included, this serves to create greater diversity in the input data to reinforce the model to generalize. Implementing regularization in the model can also help mitigate overfitting, such as "dropout"

in layers within the MLP, the Stochastic Depth (DropPath) that prevents all layers from being activated in each training step, Weight Decay with L_2 penalty to the size of the weights in the optimizer.

Given the positive results, the HAM B implementation demonstrates that complementary attention mechanisms help capture additional features to converge more quickly in this case by incorporating the inherent nonlinearity of the soft attention mechanism to capture more complex features of the input data to calculate output attention scores. Further testing and implementation of other attention mechanisms are necessary to achieve more comprehensive and satisfactory results.

Acknowledgments. This research was not funded by private entities. The work was conducted as part of graduate studies with support from the National Program for Quality Graduate Studies (PNPC) of the National Council for Humanities, Science, and Technology (CONAHCYT) in Mexico.

References

1. A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, I. Polosukhin: Attention is all you need. In: *Advances in Neural Information Processing Systems*, vol. 30, pp. 5998–6008 (2017)
2. A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly, J. Uszkoreit, N. Houlsby: An Image is Worth 16x16 Words. *Transformers for Image Recognition at Scale*, arXiv preprint arXiv:2010.11929 (2020)
3. Wang, S., Zhou, B., Jiang, J., Chen, Y.: Linformer: Self-Attention with Linear Complexity. arXiv preprint arXiv:2006.04768 (2020). <https://doi.org/10.48550/arXiv.2006.04768>
4. Beltagy, I., Peters, M.E., Cohan, A.: Longformer: The Long-Document Transformer. arXiv preprint arXiv:2004.05150 (2020). <https://doi.org/10.48550/arXiv.2004.05150>
5. Choromanski, K., Likhoshesterov, V., Dohan, D., Song, X., Gane, A., Sarlos, T., et al.: Rethinking Attention with Performers. In: *International Conference on Learning Representations (ICLR)*, arXiv preprint arXiv:2009.14794 (2021). <https://doi.org/10.48550/arXiv.2009.14794>
6. Bahdanau, D., Cho, K., Bengio, Y.: Neural machine translation by jointly learning to align and translate. arXiv preprint arXiv:1409.0473 (2015). <https://doi.org/10.48550/arXiv.1409.0473>
7. Luong, M.-T., Pham, H., Manning, C.D.: Effective approaches to attention-based neural machine translation. In: *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pp. 1412–1421 (2015). <https://doi.org/10.18653/v1/D15-1166>
8. J.-H. Kim, K.-W. On, W. Lim, J. Kim, J.-W. Ha, and B.-T. Zhang: Hadamard Product for Low-Rank Bilinear Pooling. In: *Proceedings of the 5th International Conference on Learning Representations (ICLR)* (2017) [Online]. Available: <https://openreview.net/pdf?id=r1rhWnZkg>
9. K. He, X. Zhang, S. Ren, J. Sun: Deep Residual Learning for Image Recognition. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 770–778 (2016)

10. Meta AI: PyTorch Foundation: A New Home for the Open Source AI Framework. Meta AI Blog (2022) [Online]. Available: <https://ai.meta.com/blog/pytorch-foundation/>
11. A. Krizhevsky: Learning Multiple Layers of Features from Tiny Images. Technical Report, University of Toronto (2009) [Online]. Available: <https://www.cs.toronto.edu/~kriz/learning-features-2009-TR.pdf>
12. A. Krizhevsky, I. Sutskever, G. E. Hinton: ImageNet Classification with Deep Convolutional Neural Networks. In: Advances in Neural Information Processing Systems, vol. 25, pp. 1097–1105 (2012)
13. Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., et al.: An image is worth 16x16 words: Transformers for image recognition at scale. arXiv preprint arXiv:2010.11929 (2021). <https://doi.org/10.48550/arXiv.2010.11929>

Una aproximación al marco jurídico de la Inteligencia Artificial en México

Antonio de Jesús Remes Díaz

Universidad Metropolitana de Monterrey,
México

antonio.remes@umm.edu.mx

Resumen. La regulación jurídica de la Inteligencia Artificial en México es un reto de Estado que pone a prueba a nuestro país en materia de Ciencia, Tecnología, Innovación y Desarrollo Económico. A partir de la noción de la complejidad jurídica que reviste a la Inteligencia Artificial y a la experiencia que proviene de la Unión Europea con la aprobación de su Reglamento de Inteligencia Artificial, desarrollamos un ciclo regulatorio centrado en los derechos humanos que involucra a diversas ramas jurídicas. Para robustecer nuestra exposición, presentamos el caso de las Redes Eléctricas Inteligentes en el sector energético como el antecedente dentro de nuestra legislación que consideramos más acorde con el que podemos aproximarnos de forma moderada al diseño regulatorio de la Inteligencia Artificial en el país. En ese orden de ideas y tomando en cuenta los esfuerzos legislativos durante los últimos años hacia esta disciplina, mediante una metodología similar a la utilizada en Seguridad Nacional y Defensa, una Pauta Regulatoria de la Inteligencia Artificial se propone de manera indicaria a razón de cinco requisitos: Política Pública, Diplomacia, Censo de Infraestructura Estratégica, Derechos Fundamentales y Marco Jurídico.

Palabras clave: Marco jurídico, inteligencia artificial, redes eléctricas inteligentes.

An Approach to the Legal Framework of Artificial Intelligence in Mexico

Abstract. The legal regulation of Artificial Intelligence in Mexico is a state challenge that tests our country in terms of Science, Technology, Innovation, and Economic Development. Based on the notion of the legal complexity surrounding Artificial Intelligence and the experience gained by the European Union with the approval of its Artificial Intelligence Regulation, we develop a regulatory cycle centered on human rights that involves various legal branches. To strengthen our presentation, we present the case of Smart Grids in the energy sector as the precedent within our legislation that we consider most consistent with which we can moderately approximate the regulatory design of Artificial Intelligence in the country. In this context, and taking into account the legislative efforts in recent years regarding this discipline, using a methodology similar to that used in National Security and Defense, a Regulatory Guideline for Artificial Intelligence is proposed tentatively based on five requirements: Public Policy, Diplomacy, Strategic Infrastructure Census, Fundamental Rights, and Legal Framework.

Keywords: Legal framework, artificial intelligence, smart grids.

1. Marco jurídico de la Inteligencia Artificial en México

1.1 Complejidad jurídica de la Inteligencia Artificial

La regulación jurídica de la Inteligencia Artificial en México es un reto de Estado que requiere, entre otras cuestiones, un enfoque multidisciplinario y transversal que advierta el estado del arte de la disciplina, el proceso de madurez de dicha tecnología y la asimilación de esta, a fin de comprender su interdependencia directa hacia los derechos humanos, considerando el contexto que distingue a nuestra nación dentro de la región denominada *Sur Global* [43], con el propósito de satisfacer las exigencias culturales y sociales del conocimiento de frontera.

Es importante decir que desarrollar un régimen jurídico para una de las mayores innovaciones tecnológicas que se han creado en la historia de la humanidad no es una tarea sencilla y su complejidad que representa para el Poder Legislativo en años recientes ha puesto a prueba al verdadero estado que guarda nuestro país en materia de Ciencia, Tecnología, Innovación y Desarrollo Económico, tal y como lo muestra la *Evaluación de estado de preparación de la inteligencia artificial*, elaborado por la Organización de las Naciones Unidas para la Educación, la Ciencia y la Cultura (UNESCO) en el año 2024 [43].

Considerando una década de incipientes esfuerzos provenientes del campo de la Ética de la Inteligencia Artificial a nivel internacional, principalmente por el liderazgo del *Future of Life Institute*, la organización con mayor influencia a partir de sus Cartas Abiertas, de la Conferencia de Asilomar sobre Inteligencia Artificial Beneficiosa del año 2017 y de coadyuvar con la Organización para la Cooperación y el Desarrollo Económico (OCDE), la Organización de Naciones Unidas (ONU), o bien, la propia Unión Europea, en donde contribuyó para el desarrollo del Reglamento de Inteligencia Artificial, publicado en el Diario Oficial de la Unión Europea el 12 de julio del 2024 [19, 20], es pertinente que el régimen jurídico que se proponga para nuestro país cumpla con los siguientes elementos:

- **UNO**, tener en cuenta los riesgos futuros de la inteligencia artificial;
- **DOS**, garantizar la protección de los derechos fundamentales y el Estado de Derecho, y
- **TRES**, impulsar la innovación y el desarrollo tecnológico [27].

En esas circunstancias, la complejidad jurídica que reviste a la inteligencia artificial podrá ser colmada con el ingenio característico del legislador mexicano.

Para aproximarnos a lo anterior, propongo un ciclo regulatorio que toma el enfoque de Sartor y otros en el diseño regulatorio de la Inteligencia Artificial desde su experiencia en la Unión Europea y que considera la postura de los estudios de la complejidad en esta asignatura elaborados durante la última década, como una guía temática que funciona a manera de sinergia entre la Ética, el Derecho y, por supuesto, los Derechos Humanos [4,21,26,35-39].

A fin de materializar lo anterior, resulta conveniente realizar una aproximación a lo que ha sucedido durante este proceso regulatorio en otras partes del mundo, como es el caso de Europa, en donde el pasado 12 de julio del 2024 en el Diario Oficial de la Unión Europea se publicó el *Reglamento de Inteligencia Artificial*, después de un complejo proceso de análisis y consulta entre los Estados miembros, en donde se hace notar el acervo cultural que sustenta a una disposición jurídica con sofisticada temática y confección legal a detalle [18].

En ese orden, la estructura del Reglamento de Inteligencia Artificial cuenta con una regulación a partir de aspectos fundamentales, como lo es la clasificación de la propia inteligencia artificial en función del riesgo que representa para evitar las prácticas totalmente violatorias a los derechos humanos, como la puntuación social y la manipulación del usuario [18].

Para sustentar nuestra argumentación, consideramos que es clara la intención del legislador europeo, en el recital número 48, de esta pretensión regulatoria desde el enfoque de la interdependencia jurídica por razón de la complejidad, cuando *“La magnitud de las consecuencias adversas de un sistema de IA para los derechos fundamentales protegidos por la Carta es especialmente importante a la hora de clasificar un sistema de IA como de alto riesgo. Entre dichos derechos se incluyen el derecho a la dignidad humana, el respeto de la vida privada y familiar, la protección de datos de carácter personal, la libertad de expresión y de información, la libertad de reunión y de asociación, el derecho a la no discriminación, el derecho a la educación, la protección de los consumidores, los derechos de los trabajadores, los derechos de las personas discapacitadas, la igualdad entre hombres y mujeres, los derechos de propiedad intelectual, el derecho a la tutela judicial efectiva y a un juez imparcial, los derechos de la defensa y la presunción de inocencia, y el derecho a una buena administración. Además de esos derechos, conviene poner de relieve el hecho de que los menores poseen unos derechos específicos consagrados en el artículo 24 de la Carta y en la Convención sobre los Derechos del Niño de las Naciones Unidas, que se desarrollan con más detalle en la observación general n.º 25 de la Convención sobre los Derechos del Niño de Naciones Unidas relativa a los derechos de los niños en relación con el entorno digital. Ambos instrumentos exigen que se tengan en consideración las vulnerabilidades de los menores y que se les brinde la protección y la asistencia necesarias para su bienestar. Cuando se evalúe la gravedad del perjuicio que puede ocasionar un sistema de IA, también en lo que respecta a la salud y la seguridad de las personas, también se debe tener en cuenta el derecho fundamental a un nivel elevado de protección del medio ambiente consagrado en la Carta y aplicado en las políticas de la Unión. [18]”*

Este recital es igualmente aplicable, de forma sistemática, a lo establecido en los diversos 52 y 54 a 63, del citado Reglamento, o bien, a lo previsto por los artículos 2 y 6, así como el anexo III, Sistemas de IA de alto riesgo contemplados en el apartado 2 del artículo 6 [18].

De acuerdo con el artículo 6 del Reglamento de Inteligencia Artificial, un sistema de inteligencia artificial es de alto riesgo cuando *represente un riesgo significativo de perjuicio para la salud, la seguridad o los derechos fundamentales* [18].

En ese sentido, de conformidad con lo dispuesto en el anexo III del Reglamento de Inteligencia Artificial, son sistemas de alto riesgo en esta disciplina aquellos que tocan los siguientes tópicos:

- a) Datos biométricos,
- b) Infraestructuras Críticas,
- c) Educación y formación profesional,
- d) Empleo, gestión de trabajadores y acceso al autoempleo,
- e) Acceso y disfrute de servicios privados esenciales y de los servicios públicos y prestaciones esenciales,
- f) Fuerzas y Cuerpos de Seguridad,
- g) Gestión de Migración, asilo y control de fronteras,
- h) Administración de Justicia y Procesos Democráticos [18].

Entendiendo la estructura otorgada por el legislador europeo para los sistemas de alto riesgo, podemos aproximarnos a la figura de la *evaluación de la conformidad*, reconocida en el citado Reglamento en su artículo 43, para los efectos de garantizar que los sistemas de inteligencia artificial que se consideran de alto riesgo sean fiables, antes de introducirse al mercado o ponerse en servicio [18].

La misma suerte guarda el contenido de lo dispuesto por el artículo 27 de dicha regulación europea, en materia de *evaluación de impacto sobre los derechos fundamentales*, el cual establece que los implantadores de sistemas de Inteligencia Artificial de alto riesgo identifiquen los riesgos específicos para los derechos de las personas o grupos que puedan verse afectados y determinen las medidas que deben adoptarse [18].

En ese orden de ideas, el asidero regulatorio de la Inteligencia Artificial que desarrollamos funciona como un enfoque garantista centrado en la figura de los Derechos Humanos, en concreto, en el desarrollo progresivo de los Derechos Económicos, Sociales, Culturales y Ambientales (D.E.S.C.A.), especialmente por lo que toca al Derecho de los Consumidores, al igual que la justa calificación dentro del ámbito de competencia del Derecho de la Propiedad Intelectual, del Derecho a la Protección de Datos Personales y del Derecho Internacional Humanitario, por citar una estructura esencial normativa que dinamiza nuestro ciclo regulatorio para con el resto de las ramas disciplinarias a ponderar mediante su tratamiento tradicional, como el Derecho Civil y el Derecho Mercantil, el Derecho Penal y el Derecho Administrativo, o bien, el Derecho Laboral [27].

La aproximación temática de dicho ciclo regulatorio la podemos visualizar en la siguiente Figura 1.



Fig. 1. Ciclo regulatorio de la Inteligencia Artificial a partir de los Derechos Humanos [27].

Bajo este enfoque, es posible que todas las ramas jurídicas se integren en atención a su propia especialidad, objeto de estudio y método, aproximándonos a la valoración más esencial posible dentro de las Ciencias Jurídicas, pues de cierto es que gran parte de la inquietud que reviste a la complejidad de la inteligencia artificial es abonada por la propia necesidad de diseñar un marco jurídico de la inteligencia artificial, que bien puede ser una ley, un reglamento, una reforma a algún artículo de la Constitución Política de los Estados Unidos Mexicanos, o bien, la creación de alguna Norma Oficial Mexicana o un Estándar de Competencia, considerando la existencia de algunos estándares de IEEE e ISO que involucran lo referente a la inteligencia artificial [22,23,27].

Diversas experiencias previas en nuestro Sistema Jurídico Mexicano me llevan a citar que este problema de la complejidad no es novedoso y se ha presentado en profundos procesos de Reforma Constitucional, como el caso de la Reforma del Nuevo Sistema de Justicia Penal en el año 2008, la Reforma en materia de Juicio de Amparo y Derechos Humanos del año 2011, o bien, la Reforma Laboral del año 2019 [27].

De la misma forma, algunas novedosas leyes en México expedidas entre los años 2023 y 2024, dígame la Ley General de Mecanismos Alternativos de Solución de Controversias o el Código Nacional de Procedimientos Civiles y Familiares, que regulan diversos aspectos que involucran el papel de esta disciplina mediante la Solución de Controversias en Línea, el funcionamiento del Metaverso o las Diligencias Virtuales, incluyendo la figura de Videoconferencias en Procesos Internacionales, no dejan de revestir el problema de complejidad por lo que significa la adaptación de la tecnología al sistema de justicia, recordando en este punto los casos de implementación del Juicio en Línea ante el Tribunal Federal de Justicia Administrativa, o bien, del Portal de Servicios en Línea del Poder Judicial de la Federación y el trámite del Juicio de Amparo en línea [7,10,29].

2. El caso de las redes eléctricas inteligentes en el sector energético

Con el propósito de robustecer nuestra disertación, a partir del reto que significa esta disciplina desde el punto de vista jurídico, me permito exponer la experiencia temática

del régimen jurídico de las Redes Eléctricas Inteligentes (Smart Grid) en México, como el antecedente más acorde con el que podemos ponderar, de forma moderada, el diseño regulatorio de la inteligencia artificial en el Sistema Jurídico Mexicano.

La regulación jurídica de las Redes Eléctricas Inteligentes representó un reto a nivel técnico, científico y filosófico, pues su respaldo institucional provino de parte del Gobierno de los Estados Unidos de América, tal y como consta en el documento *Marco Regulatorio de la Red Eléctrica Inteligente*, fechado en septiembre del año 2014, dirigido a la entonces Comisión Reguladora de Energía [13].

Es importante mencionar que en ese documento público se considera el papel del desarrollo de los vehículos eléctricos, de sistemas de almacenamiento de energía y un robusto capítulo de evaluaciones preliminares de desarrollo, de impacto ambiental, así como un modelo de negocios para proveedores dentro de la Red Eléctrica Inteligente [13].

En materia de Redes Eléctricas Inteligentes, el principal responsable de la implementación de esta tecnología crítica para el desarrollo nacional es la Comisión Federal de Electricidad (CFE), la cual adaptó dentro de su Programa de Ampliación y Modernización de las Redes Generales de Distribución 2015-2019, para los efectos de modernizar el Sistema Eléctrico Nacional y disminuir las pérdidas de energía, en especial, aquellas provocadas por usos ilícitos [14].

De la misma forma, entre los años 2015-2017, la entonces Empresa Productiva del Estado llevó a cabo la colocación de equipos de medición avanzada (Smart Meters) en diversas partes del país [30], los cuales, desde el punto de vista científico, como sugieren Armel y otros, representan el *santo grial* de la desagregación de energía [3, 24], aunque ello significa, desde el punto de vista jurídico, un tratamiento de datos personales a partir de la privacidad del domicilio del usuario sin que exista una orden judicial de por medio, cuestión que se encuentra en contravención a lo dispuesto por el artículo 16, décimo segundo párrafo, de la Constitución Política de los Estados Unidos Mexicanos [30].

Otros aspectos jurídicos relevantes en materia de Redes Eléctricas Inteligentes implicaban una modificación al contrato de suministro de energía eléctrica en baja tensión que proveía la entonces CFE Suministrador de Servicios Básicos, la tecnología de desagregación de energía al interior del domicilio y un nuevo marco regulatorio de tarifas eléctricas [30].

En esa época, consta por demás que la CFE desarrolló su tecnología de Redes Eléctricas Inteligentes con algunas patentes, como la diversa titulada *Sistema de monitoreo y control de medición de energía*, con el número de solicitud PCT/MX2015/000156, la cual fue reivindicada con el estado que guardaban sus componentes técnicos dentro del proceso de Distribución y Comercialización, como son los sistemas legados SIMED, SICOM y SICOSS [44].

Ahora bien, cuando se expidió la Ley de Transición Energética y la Ley de Industria Eléctrica al amparo de la Reforma Constitucional en materia de Energía del año 2013, las Redes Eléctricas Inteligentes se reconocieron en ambas leyes federales y contaban con un marco jurídico que vinculaba de forma directa a la competencia de la Procuraduría Federal del Consumidor (PROFECO) [11, 12]. En ese sentido, en el artículo 121, fracciones I a III, de la Ley de Transición Energética, se establecen multas de hasta veinte mil veces el salario mínimo a la persona física o moral que importe, distribuya o comercialice equipos o aparatos que incluyen información falsa o

incompleta, que implique engaño al consumidor, o bien, constituya una práctica que pueda inducir al error [11].

En todo caso, el escalamiento de esta tecnología quedó relegado de parte de la CFE y de las autoridades del sector, a un nivel de proyectos institucionales dentro de su Programa de Ampliación y Modernización de las Redes Generales de Distribución que no corresponden al Mercado Eléctrico Mayorista 2024-2038, sujeto a la disponibilidad de los recursos presupuestarios y teniendo como principal acto de implementación el cambio de medidores eléctricos [15].

Ahora bien, de cierto es que el nuevo régimen jurídico que se ha aprobado para la especialidad, con motivo de la Reforma Constitucional en materia de Energía publicada en el Diario Oficial de la Federación el 31 de octubre del 2024, o bien, la publicación de las nuevas Leyes Secundarias de la materia el 18 de marzo del 2025 [6,8,9], puede considerarse una regresión jurídica de quince años de esfuerzos interinstitucionales en el orden de la cooperación internacional y del desarrollo tecnológico para el país.

En ese sentido, en el texto de la nueva Ley de Planeación y Transición Energética publicada en el Diario Oficial de la Federación el 18 de marzo del 2025, en sus artículos 29, fracción VI, y 30, se establece que *el Plan para la Transición Energética y el Aprovechamiento Sustentable de la Energía (PLATEASE) debe identificar, evaluar, diseñar, innovar, establecer e instrumentar estrategias, acciones y proyectos en materia de redes eléctricas inteligentes y almacenamiento de energía* [8].

En contraste, la Ley de Transición Energética del año 2015 abrogada, establecía un marco jurídico robusto y sólido, desde una definición en el artículo 3, fracción XXXVIII, así como el Capítulo VI, del Programa de Redes Eléctricas Inteligentes, en sus artículos 37 a 42, correlacionando su aplicación a lo establecido en la Ley de Industria Eléctrica, en sus artículos 3, fracción XXXIV, 12, fracción XXXVIII y 14, fracción II [11,12].

En lo que toca al parámetro de los derechos humanos afectados con esta regresión, el artículo 38, fracción VIII, de la Ley de Transición Energética abrogada, reconocía lo referente a la *información hacia los consumidores y las opciones para el control oportuno de sus recursos* [11], uno de los elementos más indispensables en que descansa tanto el Derecho de los Consumidores, como el Derecho Ambiental, en especial, el derecho de acceso a la información ambiental.

Apelando al estado que guarda la literatura de dicha tecnología emergente, salvo la mejor opinión que este foro científico pueda tener, las Redes Eléctricas Inteligentes son una tecnología crítica crucial para el desarrollo de la Industria Energética, en donde se integran el resto de las tecnologías emergentes en materia de Energías Limpias y Energías Renovables.

Desde luego, la inteligencia artificial es la piedra angular para esta tecnología.

3. Pauta regulatoria de la Inteligencia Artificial: Propuesta

En el momento en que se propone este artículo para el honorable XVII Congreso Mexicano de Inteligencia Artificial, resulta relevante recordar que, ante el Senado de la República, las Comisiones de Ciencia y Tecnología e Innovación y de Análisis, Seguimiento y Evaluación sobre la aplicación y desarrollo de la Inteligencia Artificial

en México, han emprendido los trabajos para aproximarse a un marco jurídico acorde a las exigencias nacionales [40].

De la misma forma, ante la Cámara de Diputados, las Comisiones Unidas de Seguridad Ciudadana y de Ciencia, Tecnología e Innovación, en febrero del año 2024, aprobaron la Metodología para analizar la iniciativa que expide la Ley Federal de Ciberseguridad; sin embargo, un mes más tarde, sería retirada dicha iniciativa [5].

De acuerdo con los registros de iniciativas de ley y de Reforma a artículos de la Constitución Política de los Estados Unidos Mexicanos con los que cuenta la Cámara de Diputados desde la LXIV Legislatura, se han identificado un total de 45 proyectos en materia de Inteligencia Artificial, así como 19 en materia de Ciberseguridad. Por otro lado, en el Senado de la República, encontramos en sus registros de información y transparencia parlamentarias desde la LXIV Legislatura un total de 14 proyectos en materia de Inteligencia Artificial y 10 en materia de Ciberseguridad [5,40].

En atención a lo anterior, las respuestas a las solicitudes de acceso a la información folio 330030224000303 y folio 330030324000327 obtenidas a través de la Plataforma Nacional de Transparencia, en el Oficio No. IBD/ST/181/24, de fecha 3 de marzo del 2024, emitido por la Secretaría Técnica del Instituto Belisario Domínguez, o bien, el Oficio CEDIP/LXV/CT/90/2024, de fecha 15 de marzo del 2024, expedido por la Coordinadora Técnica del Centro de Estudios de Derecho e Investigaciones Parlamentarias, ponen de relieve la información técnica, científica, filosófica, así como empresarial, industrial y de la sociedad civil para el desarrollo de las iniciativas de ley (marco jurídico) en materia de Inteligencia Artificial, Ciberseguridad y Tecnologías Emergentes en México, en el periodo comprendido del 1 de septiembre del 2014 al 1 de marzo del 2024 [42].

Desde luego, la Sociedad Mexicana de Inteligencia Artificial (SMIA) ha estado de cerca en el contexto previamente descrito, mientras que la Academia Mexicana de Ciberseguridad y Derecho digital elaboró el documento *Panorama de la Inteligencia Artificial en México: la relevancia del Sandbox*, publicado en el mes de marzo del 2024, el cual contó con el financiamiento del Reino Unido [1].

Sin embargo, apelando al caso de las Redes Eléctricas Inteligentes que hemos expuesto y por lo que he podido constatar desde la elaboración del papel *Ética de las Máquinas y Justicia* para el XV Congreso Nacional de Abogados de la Barra Mexicana Colegio de Abogados en lo que toca al incipiente desarrollo del campo de la Ética de la Inteligencia Artificial [28, 29, 31], necesitamos moderación en estos esfuerzos legislativos que han sido emprendidos pues de cierto es que una Ley en materia de Inteligencia Artificial es solo uno de los componentes regulatorios a seguir y no resuelve por sí misma las necesidades específicas de la asignatura, considerando por demás la complicada categorización que presentan C. Huang y otros [16], o bien, los planteamientos de control y gobernanza que propone Rossi[32]. En ese contexto, el documento *AAAI 2025 Presidential Panel on the Future of AI Research*, publicado en el mes de marzo del 2025, pone de relieve una prospectiva compleja para abordar el marco regulatorio de la asignatura [2].

De cierto es que la regulación jurídica de la inteligencia artificial es una cuestión que guarda una apariencia digna del género de la ciencia ficción y muchos de los temas de debate de la asignatura, en especial aquellos que tocan lo referente de los riesgos futuros, como la Singularidad, la Superinteligencia, la Inteligencia Artificial General (IAG), los Sistemas de Monitoreo, los Sistemas de Armas Autónomos Letales

(S.A.A.L.) y la Paradoja Energética de la Inteligencia Artificial, son materia de reflexión [17, 27].

Por ende, lo más razonable es proponer una *Pauta Regulatoria de la Inteligencia Artificial* que tenga un propósito similar a lo que representan las actividades de planeación en materia de Seguridad Nacional y Defensa, como provee la metodología de la Seguridad Nacional diseñada por Santos Camaal [33, 34], recordando que uno de los principales esfuerzos a nivel internacional que podemos mencionar para nuestra asignatura se encuentran en el *Informe de la Comisión de Seguridad Nacional en materia de Inteligencia Artificial del Gobierno de los Estados Unidos de América*, publicado en el mes de marzo del 2021 [25].

En general, la Pauta Regulatoria de la Inteligencia Artificial que proponemos debe satisfacer el parámetro técnico, científico y filosófico del despliegue de esta tecnología crítica en el ámbito de la política pública institucionalizada, el desarrollo de planes, programas gubernamentales y censos temáticos en materia de Ciencia, Tecnología e Innovación, formación exclusiva de recursos humanos y desarrollo económico.

La propuesta que planteamos requiere, entre otros aspectos jurídicos, tomar en cuenta el orden de la Diplomacia y las obligaciones en materia de cooperación internacional que guarda el Gobierno de México a la luz de distintos tratados internacionales, destacando el Tratado México-Estados Unidos de América- Canadá (T-MEC) y el contenido de lo dispuesto en su artículo 19.15, *ciberseguridad*.

De la misma forma, considero que el desarrollo de esta Pauta Regulatoria de la Inteligencia Artificial necesitará un Censo de Infraestructura Estratégica en donde se establezcan los elementos formales para la asimilación de la tecnología de vanguardia atendiendo a las exigencias del estado del arte y de los recursos con los que contamos en México, en especial, bajo un parámetro de Seguridad Energética, Seguridad Hídrica y Seguridad Alimentaria.

La Pauta Regulatoria de Inteligencia Artificial exige una cosmovisión en tiempo real del impacto que representa esta tecnología crítica para los derechos humanos. Sobre este punto, puedo anticipar que el principal reto es colmar la tutela efectiva de los derechos fundamentales y su justiciabilidad, considerando la relevancia y el papel que guarda para esos fines el desarrollo de jurisprudencia temática mediante litigio estratégico de derechos humanos, así como casos contenciosos en materia de daños y responsabilidad patrimonial del Estado.

Finalmente, proponemos un marco jurídico dentro de la Pauta Regulatoria de la Inteligencia Artificial mediante una Ley de Inteligencia Artificial, en donde se establezcan parámetros similares a los reconocidos por la Unión Europea en su Reglamento de la materia, considerando en todo caso los siguientes aspectos de nuestra realidad:

- a) Identificar de manera extensiva los sistemas prohibidos de Inteligencia Artificial en el marco del Derecho Internacional Humanitario, así como los riesgos futuros que representan la Singularidad, la Superinteligencia, la Inteligencia Artificial General (IAG), los Sistemas de Armas Autónomos Letales (S.A.A.L), los Sistemas de Monitoreo y Puntuación Social, al igual que la Paradoja Energética de la Inteligencia Artificial, y
- b) Un capítulo temático al Desarrollo Tecnológico, Innovación y Financiamiento Sostenible, con reconocimiento de derechos laborales de profesores,

Tabla 1. Pauta Regulatoria de la Inteligencia Artificial.

Requisitos	Concepto	Elementos característicos
Política Pública	Programas y Planes de IA	Técnico, científico y filosófico
Diplomacia	Tratados Internacionales, T- MEC (Ciberseguridad)	Cooperación Internacional
Censo de Infraestructura Estratégica	Elementos formales para la asimilación de la IA atendiendo a las exigencias del estado del arte	Seguridad Energética, Seguridad Hídrica y Seguridad Alimentaria
Derechos fundamentales	Tutela efectiva y justiciabilidad de los derechos fundamentales	Litigio estratégico, daños y responsabilidad patrimonial del Estado
Marco jurídico	Ley de Inteligencia Artificial	Riesgos futuros de IA (Singularidad, Superinteligencia, IAG, S.A.A.L, Sistemas de Monitoreo y Paradoja Energética de IA)

científicos y expertos en la asignatura, al igual que un Programa Gubernamental al amparo del Plan Nacional de Desarrollo.

En la Tabla 1 que a continuación presentamos, podemos constatar estos cinco requisitos que integran de forma indiciaria la propuesta de este artículo.

4. Conclusión

La propuesta de Pauta Regulatoria de la Inteligencia Artificial que desarrollamos es el resultado de una aproximación temática e indiciaria al complejo asunto que significa la regulación de la inteligencia artificial, no solo en México, sino en todo el mundo, la cual consideramos que debe profundizarse en su desarrollo e investigación. Considerando el caso de las Redes Eléctricas Inteligentes en las leyes secundarias de la Reforma Energética de los años 2013 y 2024, podemos contar con un parámetro de referencia moderado para aproximarnos al marco jurídico de la Inteligencia Artificial.

En prospectiva, es indispensable un mayor desarrollo en el campo de la Ética de la Inteligencia Artificial en el país, como así lo advierten los distintos expertos que se han pronunciado en este terreno regulatorio [16, 32, 41], a fin de robustecer las alternativas a diversos problemas específicos que se encuentran considerados dentro de los cinco requisitos de la Pauta Regulatoria que proponemos, como el caso de los riesgos futuros, aunque sugiero conveniente conservar nuestro enfoque moderado para dichos efectos.

En todo caso, tal y como lo demuestra la experiencia de la Unión Europea con su Reglamento de Inteligencia Artificial, es posible que en México se expida, en algún momento futuro, una Ley de Inteligencia Artificial, la cual, como lo hemos demostrado, es tan solo un componente más bajo el enfoque que proponemos.

Agradecimientos. Agradezco a mi familia, a mis padres, a mi hermana y mis sobrinos, por su incondicional amor, apoyo y respaldo a lo largo de los años. De la misma forma, agradezco al Dr. Adán Octavio Ruíz Martínez (TEC Campus Monterrey), por su apoyo en la retroalimentación y revisión de estilo de este artículo.

Referencias

1. Academia Mexicana de Ciberseguridad y Derecho Digital, AMCID, Panorama de la Inteligencia Artificial en México: La relevancia del Sandbox, <https://www.amcid.org/page/sandboxregulatoriomexico>
2. Association for the Advancement of Artificial Intelligence, Future of AI Research 2025, <https://aaai.org/wp-content/uploads/2025/03/AAAI-2025-PresPanel-Report-FINAL.pdf>
3. Armel, K., Gupta, A., Shrimali, G., Albert, A.: Is Disaggregation the Holy Grail of Energy Efficiency? The Case of Electricity, *Energy Policy*, 52, pp. 213–234, (2013) doi: 10.1016/j.enpol.2012.08.062.
4. Bresciani, P.F., Palmirani, M.: Constitutional Opportunities and Risks of AI in the Law-Making Process. *Federalismi.it*, 2, pp.1–18 (2024)
5. Cámara de Diputados, Gaceta Parlamentaria, <https://gaceta.diputados.gob.mx/>
6. Cámara de Diputados, Constitución Política de los Estados Unidos Mexicanos, <https://www.diputados.gob.mx/LeyesBiblio/index.htm>
7. Cámara de Diputados, Código Nacional de Procedimientos Civiles y Familiares, <https://www.diputados.gob.mx/LeyesBiblio/index.htm>
8. Cámara de Diputados, Ley de Planeación y Transición Energética, <https://www.diputados.gob.mx/LeyesBiblio/index.htm>
9. Cámara de Diputados, Ley del Sector Eléctrico, <https://www.diputados.gob.mx/LeyesBiblio/index.htm>
10. Cámara de Diputados, Ley General de Mecanismos Alternativos de Solución de Controversias, <https://www.diputados.gob.mx/LeyesBiblio/index.htm>
11. Cámara de Diputados, Ley de Transición Energética (abrogada), <https://www.diputados.gob.mx/LeyesBiblio/index.htm>
12. Cámara de Diputados, Ley de Industria Eléctrica (abrogada), <https://www.diputados.gob.mx/LeyesBiblio/index.htm>
13. Comisión Reguladora de Energía CRE, Marco Regulatorio de la Red Eléctrica Inteligente, (septiembre del 2014), <https://www.gob.mx/cms/uploads/attachment/file/102193/6068.pdf>
14. Comisión Federal de Electricidad CFE, Programa General de Ampliación y Modernización de las Redes Generales de Distribución 2015-2019, <https://www.cfe.mx/distribucion/cumplimiento/pages/default.aspx>
15. Comisión Federal de Electricidad CFE, Programa General de Ampliación y Modernización de las Redes Generales de Distribución que no correspondan al Mercado Eléctrico Mayorista, 2024-2038. <https://www.cfe.mx/distribucion/cumplimiento/Documents/PAM%20RGD%202024-2038%20VF.pdf>
16. Huang, C., Zhang, Z., Mao, B. Yao, X.: An Overview of Artificial Intelligence Ethics, In *IEEE Transactions on Artificial Intelligence*, 4(4), pp. 799–819, doi: 10.1109/TAI.2022.3194503.
17. De Lucas, J.: *Blade Runner, el derecho, guardián de la diferencia*. Tirant lo Blanch México, México (2012)
18. EUR-Lex, Reglamento de Inteligencia Artificial, <https://eur-lex.europa.eu/homepage.html?lang=en&locale=es>
19. Future of Life Institute, <https://futureoflife.org/>
20. Future of Life Institute. EU Artificial Intelligence Act, <https://artificialintelligenceact.eu/es/>
21. Grochowski, M., Jablonowska, A., Lagioia, F., Sartor, G.: Algorithmic Transparency and Explainability for EU Consumer Protection: Unwrapping the Regulatory Premises. *Critical Analysis of Law*, 8(1), pp. 43–63 (2021)
22. IEEE Standard Model Process for Addressing Ethical Concerns during System Design, In: *IEEE Std 7000-2021*, pp. 1–82 doi: 10.1109/IEEESTD.2021.9536679.
23. ISO, ISO/IEC 42001:2023 Information technology — Artificial intelligence — Management system, <https://www.iso.org/es/contents/data/standard/08/12/81230.html>

24. Kolter, J.Z., Johnson, M.J.: REDD. A Public Data Set for Energy Disaggregation Research, *Artif. Intell.*, 25, (2011). <https://www.cs.cmu.edu/afs/cs.cmu.edu/Web/People/zkolter/pubs/kolter-kddsust11.pdf>
25. National Security Commission on Artificial Intelligence, The Final Report, <https://reports.nscai.gov/final-report/>
26. Reichman, A., Sartor, G.: Algorithms and Regulation, In H. W. Micklitz, O. Pollicino, A. Reichman, A. Simoncini, G. Sartor, & G. De Gregorio (Eds.), *Constitutional Challenges in the Algorithmic Society*, Cambridge University Press, cap. 8, pp. 131–181 (2021) doi: 10.1017/9781108914857.009.
27. Remes Díaz, A.: Marco jurídico de la inteligencia artificial en México: un debate entre la ética de la inteligencia artificial, los derechos humanos y el porvenir. In: 1er Congreso Internacional sobre Inteligencia Artificial y Humanidades. Universidad Autónoma de Chihuahua (2025)
28. Remes Díaz, A.: Ética de la inteligencia artificial y justicia. Número de registro 03-2024-050513592800-01, INDAUTOR (pendiente de publicación), México. (2024)
29. Remes Díaz, A.: Tratado de sabermetría legal, 2da. Edición, Editorial Amazon/KDP, México (2022)
30. Remes Díaz, A.: Tratado de derecho energético. Editorial Amazon/KDP, México (2022)
31. Remes Díaz, A.: Ética de las máquinas y justicia, In: XV Congreso Nacional de Abogados, pp. 2601–2628 (2016)
32. Rossi, F.: *Intelligenza Artificiale. Come funziona e dove ci porta la tecnologia che sta trasformando il mondo*. Laterza, Italia (2024)
33. Santos Caamal, M.: La globalización de la seguridad nacional; Centro de Estudios Superiores Navales CESNAV, México (2002)
34. Santos Caamal, M.: Metodología de la seguridad nacional. La seguridad nacional en México. Coord. José Luis Piñeyro, Universidad Autónoma Metropolitana, pp. 43–67 (2004)
35. Sartor, G.: *L'intelligenza artificiale e il diritto*. Bologna: Giappichelli, pp. 184 (2022)
36. Sartor, G.: Artificial Intelligence and Human Rights: Between Law and Ethics. *Maastricht Journal of European and Comparative Law*, 27(6), pp. 705–719 (2020) doi: 10.1177/1023263X20981566.
37. Sartor, G., Lagioia, F.: The Impact of the General Data Protection Regulation (GDPR) on Artificial Intelligence. Brussels: European Parliament, (2020). DOI:10.2861/293.
38. Sartor, G., Lagiola, F.: AI Systems Under Criminal Law: A Legal Analysis and a Regulatory Perspective. *Philosophy & Technology*, 33, pp.433–465 (2020). DOI:10.1007/s13347-019-00362-x.
39. Sartor, G.: The Way Forward for Better Regulation in the EU–Better Focus, Synergies, Data, and Technology. Bruxelles: European Parliament's Policy Department for Citizens' Rights and Constitutional Affairs. pp. 33 (2022) doi: 10.2861/183840.
40. Senado de la República, Información Parlamentaria, https://senado.gob.mx/66/informacion_parlamentaria/
41. Smuha, N.A. (ed.). *The Cambridge Handbook of the Law, Ethics and Policy of Artificial Intelligence*. Cambridge: Cambridge University Press (Cambridge Law Handbooks) (2025)
42. Transparencia para el Pueblo. Plataforma Nacional de Transparencia (PNT), respuestas a las solicitudes de acceso a la información 330030224000303 y 330030324000327, <https://www.plataformadetransparencia.org.mx/Inicio>
43. UNESCO. Global AI Ethics and Governance Observatory, México. <https://www.unesco.org/ethics-ai/es/mexico>
44. Sistema de monitoreo y control de medición de energía. WIPO-PATENTSCOPE. Número de publicación WO/2016/089194

Evaluación preliminar de reconocimiento de patrones en gotas secas de metotrexato

Carlos A. Martínez-Miwa^{1,2}, R. Magdalena Sánchez-Albores³,
Yohana J. P. Carreón³, Jorge Gonzalez-Gutiérrez³, Mario Castelán¹

¹ Instituto Politécnico Nacional,
Centro de Investigación y de Estudios Avanzados,
Mexico

² Universidad Tecnológica de Coahuila,
Mexico

³ Universidad Autónoma de Chiapas,
Mexico

{carlos.miwa, mario.castelan}@cinvestav.edu.mx,
{magdalena.sanchez, yojana.carreon, jorge.ggutierrez}@unach.mx,
carlos.miwa@utc.edu.mx

Resumen. El metotrexato (MTX) es un agente terapéutico ampliamente utilizado cuyo control de calidad es crucial para garantizar su eficacia y seguridad. En este trabajo, se propone un enfoque basado en redes neuronales convolucionales (CNN), específicamente la arquitectura VGG-16, para la clasificación automática de residuos de gotas secas de MTX con distintos grados de adulteración. Se analizaron imágenes de gotas secas de metotrexato en concentraciones del 100 %, 80 %, 60 % y 40 %. Las soluciones con 80 %, 60 % y 40 % de metotrexato correspondían a mezclas adulteradas con 20 %, 40 % y 60 % de agua, respectivamente. En todas las muestras, la concentración de NaCl se mantuvo constante en 2 %. Se implementó un método de segmentación basado en la transformada de Hough para la extracción de parches. Se generó un conjunto de datos enriquecido con imágenes de la periferia y el centro de cada gota. Los resultados obtenidos muestran una precisión promedio de clasificación superior al 90 %, con una mayor efectividad en la identificación de adulteraciones en la región periférica de las gotas. Los resultados son preliminares y muestran el potencial del aprendizaje profundo para la detección automatizada de alteraciones en MTX, con impacto en control de calidad en la industria farmacéutica.

Palabras clave: Metotrexato, gotas secas, aprendizaje profundo.

Preliminary Evaluation of Pattern Recognition in Dried Methotrexate Drops

Abstract. Methotrexate (MTX) is a widely used therapeutic agent whose quality control is crucial to ensure its efficacy and safety. This study proposes an approach based on convolutional neural networks (CNN), specifically the VGG-16 architecture, for the automatic classification of dried drop residues of MTX with varying degrees of adulteration. Images of dried methotrexate drops were analyzed at concentrations of 100 %, 80 %, 60 %, and 40 %. The 80 %, 60 %, and 40 % methotrexate solutions corresponded to mixtures adulterated with 20 %, 40 %, and 60 % water, respectively. In all samples, the NaCl concentration was kept constant at 2 %. A segmentation method based on the Hough transform was implemented for patch extraction. An enriched dataset was generated with images from both the periphery and the center of each drop. The results show an average classification accuracy above 90 %, with greater effectiveness in detecting adulteration in the peripheral region of the drops. These preliminary results highlight the potential of deep learning for the automated detection of alterations in MTX, with implications for quality control in the pharmaceutical industry.

Keywords: Methotrexate, dried drops, deep learning.

1. Introducción

El metotrexato, análogo de la ametopterina, es un agente quimioterapéutico y supresor del sistema inmunitario, empleado principalmente para tratar padecimientos oncológicos, enfermedades autoinmunes y embarazos ectópicos [15,2]. Este tipo de compuestos, desafortunadamente, es susceptible a ser adulterado, menguando su funcionalidad y comprometiendo la salud de los pacientes [6,7]. En consecuencia, la identificación de alteraciones en su composición resulta crucial [11,24,21].

Dentro de las técnicas clásicas para el control de la calidad de medicamentos se encuentran la espectroscopía de masas (MS) [20,23], la cromatografía líquida de alta resolución (HPLC) [12,8] y la cromatografía en capa fina (TLC) [18,13]. Estas estrategias posibilitan la cuantificación de contaminantes e impurezas, incluso en muy bajas concentraciones [21].

En épocas más recientes, un método ampliamente utilizado y prometedor para detectar estas alteraciones es el análisis de los patrones formados por gotas secas, los cuales reflejan las propiedades fisicoquímicas del medicamento [4]. Estos patrones se ven influenciados por diversos factores, por ejemplo, la composición química, las condiciones ambientales y los flujos internos generados durante la evaporación, como los flujos de Marangoni [16]. El efecto Marangoni es ocasionado por el gradiente de tensión superficial en la interfase de fluidos. Este fenómeno, inicialmente documentado por el científico James Thomson en 1855 [10], se manifiesta por la tendencia de las regiones de alta tensión superficial a atraer partículas

circundantes con mayor fuerza que las regiones de baja tensión superficial [9,1]. No obstante, la variabilidad en las estructuras resultantes dificulta su análisis con métodos convencionales como la matriz de co-ocurrencia de niveles de gris (GLCM) [14,5], que resulta insuficiente ante la complejidad de estas formaciones. Más aún, la inspección visual que suponen estos procesos es dependiente de la experiencia del personal a cargo y propensa a errores de medición, consecuencia de los límites de la naturaleza humana, a saber, la fatiga visual o el cansancio.

Para superar estas limitaciones, en la literatura se ha propuesto la implementación de herramientas de aprendizaje de máquina (*Machine Learning* o ML) o aprendizaje profundo (*Deep Learning* o DL). En contraste con la interpretación humana, que depende meramente de la memoria y sus capacidades físicas, las técnicas de aprendizaje de máquina “aprenden” información de bases de datos extensas. Asimismo, la información adquirida está presente casi en la totalidad del tiempo, dado que los modelos almacenan internamente los patrones aprendidos, pudiendo acceder a ella de forma automática y constante. En [22], los autores exploran las propiedades moleculares y posibles impurezas en fármacos. Dentro del estudio [19], se implementa ML para predecir efectos secundarios de diversos medicamentos, permitiendo revelar impurezas o interacciones inesperadas con otras proteínas. Por otro lado, [3] presentan un modelo de DL para analizar la actividad de moléculas de compuestos en la industria farmacéutica. Aunque ha demostrado ser una herramienta poderosa, los datos que requiere este algoritmo son privados y no están disponibles de manera pública, lo que limita su efectividad y reproducibilidad.

En este trabajo se propone el uso de aprendizaje profundo mediante redes neuronales convolucionales (CNN), específicamente la arquitectura VGG-16 [17]. Este enfoque permite una clasificación precisa de las gotas secas adulteradas de MTX, optimizando la detección de irregularidades y fortaleciendo los procesos de control de calidad en la industria farmacéutica.

Nuestros hallazgos no solo destacan el potencial del aprendizaje profundo para superar los desafíos de clasificación, sino que también proporcionan una base para implementar esta tecnología en aplicaciones de control de calidad más amplias. Este enfoque promete optimizar la detección de adulteraciones en productos farmacéuticos, garantizando estándares de calidad más altos y contribuyendo significativamente a la seguridad del paciente.

2. Análisis de imágenes de metotrexato

La base de datos analizada en el presente estudio representa imágenes de MTX bajo tres diferentes grados de adulteración (GA): MTX+2 % agua, MTX+4 % agua y MTX+6 % agua. Además, se incluye un grupo control correspondiente a MTX+10 % agua. Inicialmente, se observaron variaciones en la cantidad de imágenes por cada grado de adulteración. Se contaba con 128 imágenes de control, 65 de MTX+2 % agua, 53 para MTX+4 % agua y 48 para MTX+6 % agua. Con el propósito de garantizar la coherencia y minimizar posibles sesgos,

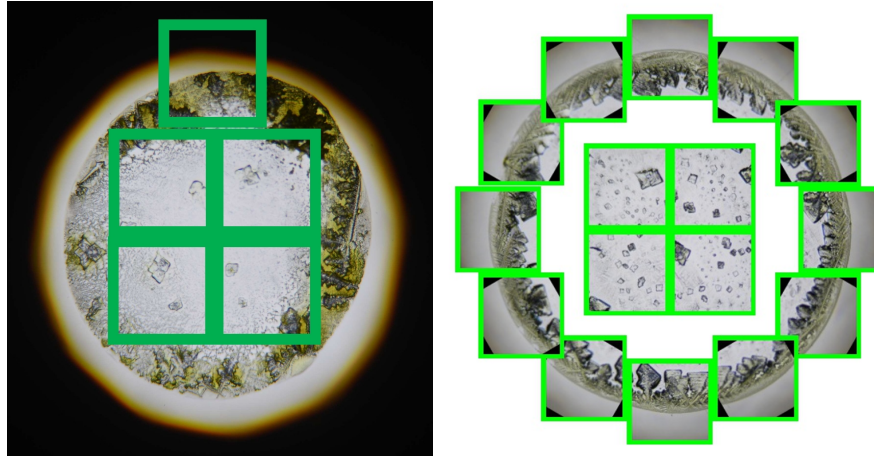


Fig. 1. Estrategia de extracción de parches propuesta para cada imagen de gotas. Se muestrearon doce parches a lo largo de la circunferencia de la gota a intervalos de 30 grados, mientras que se extrajeron ocho parches adicionales de la región central (cuatro a 0 grados y cuatro a 30 grados).

todos los conjuntos se estandarizaron a 48 imágenes, coincidiendo con el grupo más pequeño.

El proceso de selección utilizó una transformada de Hough para evaluar la imagen de cada gota, conservando sólo aquellas que mejor se ajustaban a una forma circular. Por cada imagen, se extrajeron veinte parches o regiones para ampliar los conjuntos de datos y enriquecer la información disponible para el entrenamiento. Doce parches se obtuvieron a lo largo del perímetro de cada gota a intervalos de 30°, mientras que los ocho parches restantes provinieron de cuatro secciones del centro de la gota, muestreadas a 0° y 30°. La elección de rangos de 30 grados responde a la intención de explorar diversas zonas de la imagen que proporcionen datos relevantes, evitando la superposición entre ellas y, por ende, la duplicación de información. Este procedimiento se ilustra en la figura 1. A la izquierda se aprecia la localización de las regiones a extraer, cuatro centrales y una en la periferia como ejemplo. En la parte derecha se incluye una demostración de las veinte regiones analizadas en cada gota de MTX. Como resultado, cada GA se conformó por 960 imágenes en total.

3. Experimentación

La distribución de imágenes para la etapa de experimentación se detalla en el cuadro 1. El conjunto de entrenamiento involucra 784 imágenes (70 %), 224 para prueba (20 %) y 112 de validación (10 %). Estos tres subconjuntos se evaluaron utilizando la red VGG-16 en una computadora equipada con un procesador AMD Ryzen 7, una unidad SSD de 1 TB y una GPU NVIDIA GeForce RTX 3050.

Tabla 1. Distribución de imágenes para cada GA. Algunas cantidades reflejan números decimales. Éstos corresponden a porciones de una imagen adicional. Por ejemplo, en el conjunto de entrenamiento se tienen la totalidad de parches de 33 imágenes, más el 60 % de parches de la imagen 34.

	Distribución de imágenes por GA			
	Entrenamiento	Validación	Prueba	Total
Número de Imágenes	33.6	4.8	9.6	48
Número de Parches	672	96	192	960

Tabla 2. Porcentajes medios de precisión (mAP) obtenidos tras la fase de entrenamiento de cada conjunto de datos de gotas con adulteración (GA).

Grado de Adulteración (GA)	Precisión Promedio en Entrenamiento (mAP)
MTX+2 % agua	0.98
MTX+4 % agua	0.99
MTX+6 % agua	0.91

El proceso de entrenamiento consistió en 30 épocas con imágenes redimensionadas a 224×224 píxeles. Se llevaron a cabo tres experimentos principales: 1) entrenamiento de la red VGG-16 por separado para cada GA; 2) clasificación de los parches del conjunto de prueba para determinar si pertenecían a una muestra adulterada o de control; y 3) distinción de las regiones según procedieran de la periferia o del centro de la gota.

La red de clasificación VGG-16 se entrenó en tres ocasiones, una por cada GA, para diferenciar si los parches del conjunto de entrenamiento pertenecían a una gota adulterada de MTX o a una imagen de control. Se empleó la misma cantidad de imágenes de entrenamiento para los distintos GA y el grupo control, según la tabla 1.

4. Resultados

En la tabla 2 se presentan los resultados de entrenamiento con respecto a la media de la precisión promedio (mAP, por sus siglas en inglés). Cabe destacar que, si bien la precisión disminuye a medida que aumenta el nivel de adulteración, ésta sigue siendo notablemente alta, lo que sugiere que la red aprendió efectivamente suficientes características para distinguir entre las gotas adulteradas y las de control.

Posteriormente, el conjunto de prueba de cada GA se evaluó utilizando los modelos VGG previamente entrenados. De esta manera, fue posible comprobar la efectividad del aprendizaje de la red al presentársele imágenes inéditas, es

Tabla 3. Número de parches clasificados como verdadero control (VC), falso control (FC), verdadero experimental (VE) y falso experimental (FE) para todas las GA. Para todas las bases de datos, se incluye la precisión promedio de clasificación.

Grado de Adulteración (GA)	Parches VC	Parches FC	Parches VE	Parches FE	mAP
MTX+2 % agua	190	16	176	2	0.96
MTX+4 % agua	184	11	181	8	0.95
MTX+6 % agua	188	22	170	4	0.94

decir, no incluidas en la etapa de entrenamiento. De acuerdo a las predicciones, los parches de prueba se etiquetaron como verdadero control (VC), falso control (FC), verdadero experimental (VE) y falso experimental (FE), tal y como se muestra en la tabla 3. En la última columna se indican las precisiones promedio (mAP). Es evidente el predominio de las clasificaciones verdaderas sobre las falsas, lo cual se reafirma con el alto grado de precisión de clasificación en cada una de las bases de datos. Esto pone de manifiesto la eficacia del proceso de aprendizaje conseguido con el trabajo aquí presentado.

Las imágenes de cada categoría se dividieron, a su vez, en dos subgrupos: las del centro de la gota y las de la corona. La tabla 4 muestra la precisión media de la clasificación en las tres GA. En particular, los parches extraídos de las gotas experimentales obtuvieron una mayor precisión que los de las gotas de control. Estos resultados sugieren que la corona de la gota proporciona la información más confiable para identificar la presencia de adulteración. No obstante, es importante señalar que la región central también constituye una valiosa fuente de datos, especialmente cuando la información procedente del “anillo de café”, como se conoce a la corona, es insuficiente.

La figura 2 ejemplifica la categorización de los parches de tres gotas evaporadas junto a una imagen de control, cada una de las cuales representa un grado diferente de adulteración. Los recuadros rojos indican las regiones clasificadas como falsas (FC y FE), mientras que las áreas no enmarcadas se refieren a las zonas verdaderas (VC o VE).

5. Conclusiones

Los resultados obtenidos en el presente trabajo sugieren evidencia preliminar de que el aprendizaje profundo constituye un método confiable para clasificar gotas secas adulteradas de compuestos farmacológicos, en específico metotrexato. Nuestro método ofrece una alternativa robusta a los métodos tradicionales que enfrentan desafíos significativos, tales como las limitaciones humanas o la dependencia de bases de datos sumamente extensas o de difícil acceso. La implementación de algoritmos de CNN, en especial la red VGG-16, exhibió un desempeño altamente favorable al identificar características en las imágenes de gotas

Tabla 4. Precisión de clasificación de los parches del centro y de la corona. Para cada GA, los parches se agrupan en imágenes de control (arriba) e imágenes experimentales (abajo). En general, el mejor rendimiento se obtuvo con las imágenes de la periferia, mientras que las imágenes adulteradas arrojaron la mayor precisión de clasificación.

GA	Categoría	Centro	Corona	Media
MTX+2 % agua	Control	0.87	0.96	0.92
MTX+4 % agua	Control	0.98	0.92	0.95
MTX+6 % agua	Control	0.80	0.96	0.88
Media	Control	0.88	0.95	0.92
GA	Categoría	Centro	Corona	Media
MTX+2 % agua	Experimental	0.97	1.0	0.99
MTX+4 % agua	Experimental	0.90	1.0	0.95
MTX+6 % agua	Experimental	0.97	0.98	0.98
Media	Control	0.95	0.99	0.97

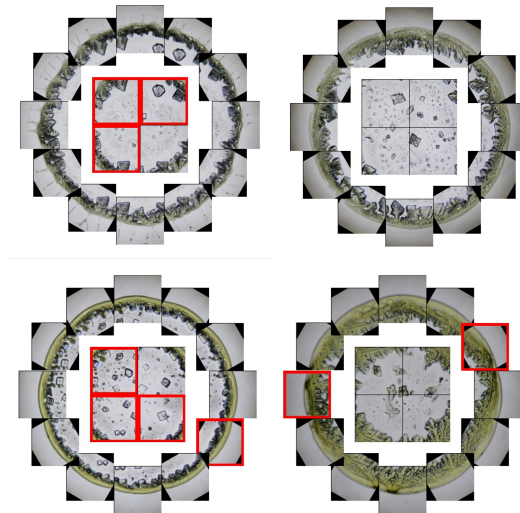


Fig. 2. Ejemplos de clasificaciones de parches para cada GA: 2 % (arriba a la izquierda), 4 % (arriba a la derecha), 6 % (abajo a la izquierda) y una imagen de control (abajo a la derecha). Las regiones rojas indican parches clasificados erróneamente (FC o FE), mientras que los parches sin colorear representan clasificaciones verdaderas (FC y FE).

secas de MTX, las cuales hicieron posible diferenciar de manera correcta aquellas imágenes de MTX con adulteración de las que no se encontraban alteradas.

Agradecimientos. Se agradece al proyecto CONAHCYT CF-2023-G-454 por el apoyo otorgado a esta investigación. C.A.M.M. reconoce a la Universidad Tecnológica de Coahuila por las facilidades y el respaldo proporcionados.

Contribución de los autores Diseño de algoritmos, programación, escritura de borrador: C.A.M.M., M.C.; generación de base de datos de MTX: R.M.S.A., Y.J.P.C., J.G.G.; pruebas de laboratorio: R.M.S.A., Y.J.P.C., J.G.G.; supervisión: J.G.G., M.C.; financiamiento: J.G.G., M.C. Todos los autores participaron en la escritura del artículo y dieron su autorización en la versión final.

Referencias

1. Abbas, M., Khan, N.: Numerical study for bioconvection in marangoni convective flow of cross nanofluid with convective boundary conditions. *Advances in Mechanical Engineering*, vol. 15, no. 11, pp. 16878132231207625 (2023)
2. American Society of Health-System Pharmacists: Methotrexate. <https://www.drugs.com/monograph/methotrexate.html> (2016), [Online; accessed 11-March-2025]
3. Cao, A., Zhang, L., Bu, Y., Sun, D.: Machine learning prediction of on/off target-driven clinical adverse events. *Pharmaceutical Research*, vol. 41, no. 8, pp. 1649–1658 (2024)
4. Carreón, Y. J., Díaz-Hernández, O., Escalera Santos, G. J., Cipriano-Urbano, I., Solorio-Ordaz, F. J., González-Gutiérrez, J., Zenit, R.: Texture analysis of dried droplets for the quality control of medicines. *Sensors*, vol. 21, no. 12, pp. 4048 (2021)
5. Carreón, Y. J., Ríos-Ramírez, M., Moctezuma, R., González-Gutiérrez, J.: Texture analysis of protein deposits produced by droplet evaporation. *Scientific reports*, vol. 8, no. 1, pp. 9580 (2018)
6. Cronstein, B. N.: The mechanism of action of methotrexate. *Rheumatic disease clinics of North America*, vol. 23, no. 4, pp. 739–755 (1997)
7. Friedman, B., Cronstein, B.: Methotrexate mechanism in treatment of rheumatoid arthritis. *Joint bone spine*, vol. 86, no. 3, pp. 301–307 (2019)
8. Jia, H., Li, R., Li, Y., Lu, F., Ma, L., Xu, X.: Improved analysis hplc-esi/triple method for mapping the methotrexate by mass spectrometry. *Journal of Chromatography B*, pp. 124529 (2025)
9. Khan, M. I., Qayyum, S., Chu, Y.-M., Khan, N. B., Kadry, S.: Transportation of marangoni convection and irregular heat source in entropy optimized dissipative flow. *International Communications in Heat and Mass Transfer*, vol. 120, pp. 105031 (2021) doi: 10.1016/j.icheatmasstransfer.2020.105031
10. Marangoni, C.: Sull'espansione delle gocce d'un liquido galleggianti sulla superficie di altro liquido. *Fratelli Fusi* (1865)
11. Megía, M. F., ALMIÑANA, M. A., CASTERA, M. T.: Manejo de la intoxicación por metotrexato: a propósito de un caso. *FARMACIA HOSPITALARIA*, vol. 2004, no. 28/5, pp. 371 (2004)
12. Muchakayala, S. K., Katari, N. K., Saripella, K. K., Schaaf, H., Marisetti, V. M., Ettaboina, S. K., Rekulapally, V. K.: Implementation of analytical quality by design and green chemistry principles to develop an ultra-high performance liquid chromatography method for the determination of fluocinolone acetonide impurities from its drug substance and topical oil formulations. *Journal of Chromatography a*, vol. 1679, pp. 463380 (2022)

13. Nizomiddinovich, T. F., Abdimannonovich, I. S., Zoirovich, A. J.: Of organic substances by thin layer chromatographic method. *TADQIQOTLAR*, vol. 57, no. 1, pp. 19–21 (2025)
14. Pal, A., Gope, A.: Texture identification in liquid crystal-protein droplets using evaporative drying, generalized additive modeling, and k-means clustering. *The European Physical Journal E*, vol. 47, no. 5, pp. 35 (2024)
15. Puig, L.: Methotrexate: New therapeutic approaches. *Actas Dermo-Sifiliográficas (English Edition)*, vol. 105, no. 6, pp. 583–589 (2014) doi: 10.1016/j.adengl.2014.05.011
16. Sefiane, K.: Patterns from drying drops. *Advances in colloid and interface science*, vol. 206, pp. 372–381 (2014)
17. Simonyan, K., Zisserman, A.: Very deep convolutional networks for large-scale image recognition. *CoRR*, vol. abs/1409.1556 (2014)
18. Sowers, M. E., Ambrose, R., Bethea, E., Harmon, C., Jenkins, D.: Quantitative thin layer chromatography for the determination of medroxyprogesterone acetate using a smartphone and open-source image analysis. *Journal of Chromatography A*, vol. 1669, pp. 462942 (2022)
19. Staszak, M., Staszak, K., Wieszczycka, K., Bajek, A., Roszkowski, K., Tylkowski, B.: Machine learning in drug design: Use of artificial intelligence to explore the chemical structure–biological activity relationship. *Wiley Interdisciplinary Reviews: Computational Molecular Science*, vol. 12, no. 2, pp. e1568 (2022)
20. Taupik, M., Djuwarno, E., Mustapa, M. A., Kunusa, W. R., Kilo, J., Sahumena, M. H.: The type fragmentation patterns confirmed acetaminophen by using liquid chromatography-mass spectroscopy (lcms) from herbal medicine (jamu). *Journal of Islamic Science and Technology*, vol. 7, no. 2, pp. 341–353 (2021)
21. Torregiani, L., Caro, Y. S., De Zan, M. M.: Monitorización terapéutica de metotrexato en pacientes leucémicos mediante cromatografía líquida de alta resolución. desarrollo, validación y aplicación clínica del método analítico. *Revista del Laboratorio Clínico*, vol. 11, no. 2, pp. 64–72 (2018)
22. Wallach, I., Dzamba, M., Heifets, A.: Atomnet: a deep convolutional neural network for bioactivity prediction in structure-based drug discovery. *arXiv preprint arXiv:1510.02855*, (2015)
23. Wang, Z., Hughes, N.: Simultaneous analysis of 7 antibiotics in urine by liquid chromatography triple quadrupole tandem mass spectroscopy and solid phase extraction. *New Zealand Journal of Medical Laboratory Science*, vol. 78, no. 3 (2024)
24. Zhao, S. S., Bukar, N., Toulouse, J. L., Pelechacz, D., Robitaille, R., Pelletier, J. N., Masson, J.-F.: Miniature multi-channel spr instrument for methotrexate monitoring in clinical samples. *Biosensors and Bioelectronics*, vol. 64, pp. 664–670 (2015) doi: 10.1016/j.bios.2014.09.082

Environmental Assessment in Pear Cactus Cultivation Using IoT + Edge AI

Luis Torres-Treviño¹, Alejandro Luna-Maldonado², Erik Palacios-Garza¹,
Erick Ordaz-Rivas¹

¹ Universidad Autónoma de Nuevo Leon,
Mexico

² ABC Institute, Rupert-Karls-University Heidelberg,
Germany

`{abc,lncs}@uni-heidelberg.de`

Abstract. A prototype based on Internet of Things (IoT) technology is presented, featuring an embedded fuzzy system that enables real-time evaluation of the condition of a nopal crop. The study analyzes the required environmental and soil conditions for optimal growth. Although the work focuses on assessing these conditions, it also considers the plant's overall health status. The presented experiments demonstrate the feasibility of evaluating nopal crops. The proposed module can be installed on a mobile robot, allowing for plant-by-plant inspection.

Keywords: IoT, edge AI, fuzzy systems, digital agriculture.

1 Introduction

1.1 Cactus Pear Uses

Cactus pear (nopal) is a versatile plant with various applications in different fields. Below are some of its most common uses [7]:

- **Culinary Uses:** Nopal pads (cladodes) and prickly pear fruits are widely consumed. They can be eaten fresh, cooked, or juiced and are commonly used in salads, stews, jams, and beverages due to their nutritional benefits.
- **Medicinal Uses:** Nopal has several potential health benefits, including:
 - Lowering blood sugar levels.
 - Reducing cholesterol.
 - Aiding digestion.
 - Possessing anti-inflammatory properties.
- **Cosmetic Uses:** Extracts from nopal are incorporated into skincare products for their hydrating, anti-aging, and soothing effects on irritated skin.
- **Animal Feed:** Nopal pads serve as a nutritious fodder option for livestock, particularly in arid regions where other feed sources are scarce.

- **Industrial Uses:** The mucilage from nopal has various industrial applications, including:
 - Water purification.
 - Bioplastics production.
 - Use as a natural thickener.
- **Environmental Uses:** Nopal plays a role in soil conservation, combating desertification, and carbon sequestration due to its resilience in harsh climatic conditions.

1.2 Factors which Affect the Crop Quality of Cactus Pear

Several factors influence the crop quality of cactus pear (*Opuntia* spp.), commonly known as nopal. Environmental conditions such as temperature, rainfall, and soil type play a crucial role, as nopal thrives in arid and semi-arid regions with well-drained, sandy-loam soils. Water availability is particularly important since drought stress can reduce pad and fruit quality, while excessive moisture may lead to fungal diseases. Proper nutrient management, especially adequate levels of nitrogen, phosphorus, and potassium, enhances growth and fruit development. Additionally, pests and diseases, including cochineal insects and fungal infections like *Phyllosticta* and *Colletotrichum*, can significantly impact yield and quality. Cultivation practices such as pruning, spacing, and harvesting techniques also contribute to the final product's texture, size, and nutritional value. Lastly, genetic factors and the choice of cultivar determine resistance to stress conditions and overall market appeal.

1.3 Uses of IoT for Digital Agriculture

The Internet of Things (IoT) is revolutionizing digital agriculture by enabling real-time monitoring, automation, and data-driven decision-making. IoT-powered smart sensors track soil moisture, temperature, and nutrient levels, allowing farmers to optimize irrigation and fertilization, reducing waste and improving crop yields. Automated irrigation systems adjust water distribution based on weather and soil conditions, conserving water while ensuring optimal plant growth. In livestock management, IoT wearables monitor animal health, movement, and feeding patterns, enhancing productivity and disease prevention. Drones and smart machinery automate planting, spraying, and harvesting, increasing efficiency and reducing labor costs. Additionally, IoT-enabled supply chain tracking ensures transparency in food production, while AI-powered predictive analytics help farmers make better decisions based on climate and market trends. By integrating IoT, agriculture becomes more sustainable, efficient, and profitable, addressing global food security challenges while minimizing environmental impact [8, 10, 11, 6, 5].

1.4 Motivations

Crop quality evaluation is a critical aspect of modern agriculture, ensuring that produce meets industry standards and consumer expectations. This process

involves assessing various factors such as plant health, growth rate, nutrient levels, and resistance to pests or diseases. Advanced technologies, including remote sensing, IoT-based monitoring, and machine learning algorithms, enhance precision in evaluating crop conditions in real time. Traditional methods, such as visual inspection and laboratory testing, are also used to analyze soil composition, moisture content, and overall plant vitality. By implementing efficient crop quality evaluation techniques, farmers can optimize yield, reduce losses, and maintain sustainable agricultural practices.

2 Hierarchical Fuzzy System for Crop Quality Evaluation

Fuzzy systems can be effectively used to analyze crops, particularly cactus pear. Since much of the knowledge related to crop evaluation is linguistic, fuzzy rules provide an ideal way to represent this information. These rules incorporate data from various sensors, including temperature, humidity, and soil characteristics, to assess and qualify the environmental conditions of the crop. This approach enables a more precise and adaptable evaluation, improving decision-making in agricultural management [4, 3, 9, 1, 2].

A Hierarchical Fuzzy System (HFS) for crop quality evaluation can be mathematically defined as a multi-level fuzzy inference system, where the output of one fuzzy system serves as an input for another. This hierarchical structure reduces complexity and improves interpretability by separating the evaluation process into smaller, manageable sub-systems cite.

2.1 Mathematical Definition

Input Variables Let $X = \{x_1, x_2, \dots, x_n\}$ be the set of input variables representing crop attributes. Each input x_i is fuzzified using membership functions $\mu_{A_i}(x_i)$, where A_i is a fuzzy set.

Fuzzy Rules At each level l , the fuzzy rules are defined as:

$$R_k^{(l)} : \text{IF } x_1 \text{ is } A_{k1}^{(l)} \text{ AND } x_2 \text{ is } A_{k2}^{(l)} \text{ AND } \dots \text{ AND } x_n \text{ is } A_{kn}^{(l)} \text{ THEN } y_k^{(l)} \text{ is } B_k^{(l)}, \quad (1)$$

where:

- $R_k^{(l)}$ is the k -th rule at level l ,
- $A_{ki}^{(l)}$ are fuzzy sets for input x_i at level l ,
- $y_k^{(l)}$ is the output of the k -th rule at level l ,
- $B_k^{(l)}$ is the fuzzy set for the output at level l .

Hierarchical Structure The output of level l becomes the input for level $l+1$. The final output Y is obtained at the last level L :

$$Y = f^{(L)}(f^{(L-1)}(\dots f^{(1)}(X))),$$

where $f^{(l)}$ is the fuzzy inference function at level l .

Defuzzification The final output Y is defuzzified using the centroid method:

$$Y_{\text{final}} = \frac{\int y \cdot \mu_Y(y) dy}{\int \mu_Y(y) dy},$$

where $\mu_Y(y)$ is the membership function of the final output.

2.2 Hierarchical Architecture

A hierarchical fuzzy system that evaluates the conditions of a nopal crop, an embedded fuzzy system is implemented within an Internet of Things (IoT) module. The architecture is shown in Figure 1. Five signals are considered to evaluate the quality of the environment (soil and atmosphere) in which the nopal crop is grown. The rules are designed taking these aspects into account, as the plant is more resistant to droughts and high temperatures, as mentioned above. To avoid the combinatorial explosion of rules, which would require including $5 \times 5 \times 5 \times 5 \times 5 = 3125$ rules, the fuzzy system with five inputs is transformed into two fuzzy systems: one with two inputs to determine the soil quality of the crop, and another with three inputs to evaluate the environment where the crop is located. The total number of rules is drastically reduced to $5 \times 5 \times 5 + 5 \times 5 = 150$ rules in total. A hierarchical fuzzy system is shown in Figure 5.

This architecture simplifies rule generation and significantly reduces the number of rules required. For instance, a rule to determine soil quality considers moisture and pH levels. In nopal (cactus pear) crops, an intermediate pH combined with low to medium moisture is ideal for enhancing soil quality. Similarly, a high or very high-quality crop environment is generally preferred for optimal growth. However, nopal crops can demonstrate resilience even in environments with medium or low quality.

3 Experimentation and Results

3.1 First Simulation: 24-Hour Evaluation

The first simulation involves a 24-hour evaluation, considering the average weather conditions in the metropolitan area of Nuevo León. The evaluation takes into account the type of soil present, which is exposed to various environmental conditions such as rain and wind, affecting its moisture and pH composition, although the latter changes only slightly.

3.2 Nopal's Resilience and Evaluation Results

The nopal is a highly drought-resistant plant, capable of withstanding extreme conditions. It is expected that under very low light conditions, the cultivation conditions are very poor; the evaluation system reflects this. Under better lighting conditions, these conditions improve significantly. However, even then, it barely surpasses a rating of 7 out of 10. Figure 3 shows the performance of the system evaluation.

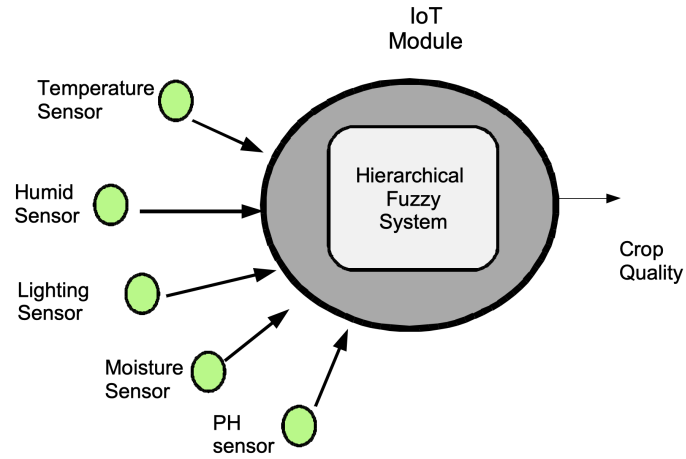


Fig. 1. IoT Module for evaluation of environment and soils of crops.

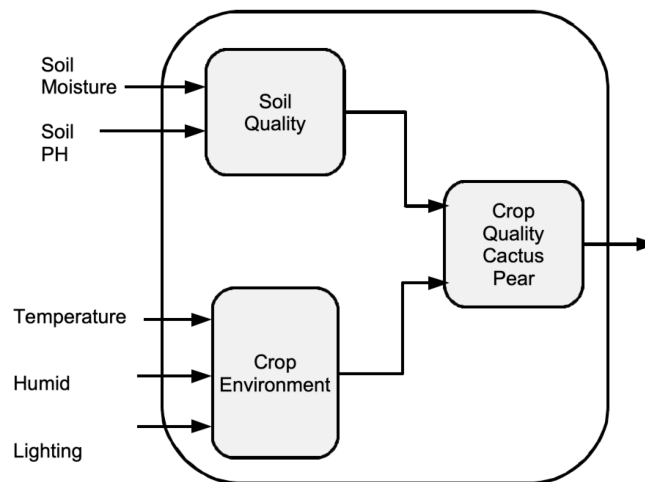


Fig. 2. Evaluation system based on a hierarchical fuzzy system.

3.3 Optimization Using the Evonorm Algorithm

Using an optimization algorithm based on the Evonorm algorithm [12], the ideal conditions to achieve a rating of 87 are as follows:

- Temperature: 21.79
- Humidity: 44
- Light intensity: 40233

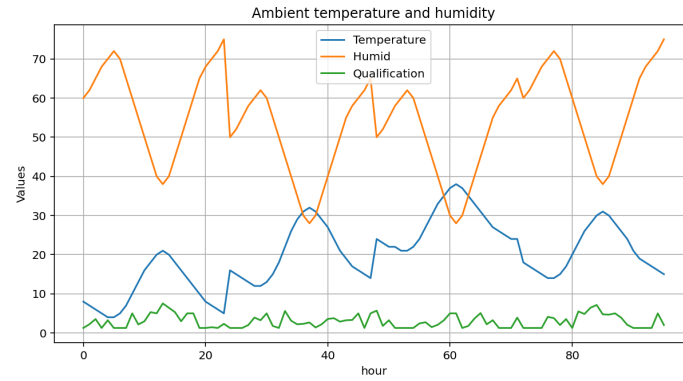


Fig. 3. Evaluation system based on a hierarchical fuzzy system.

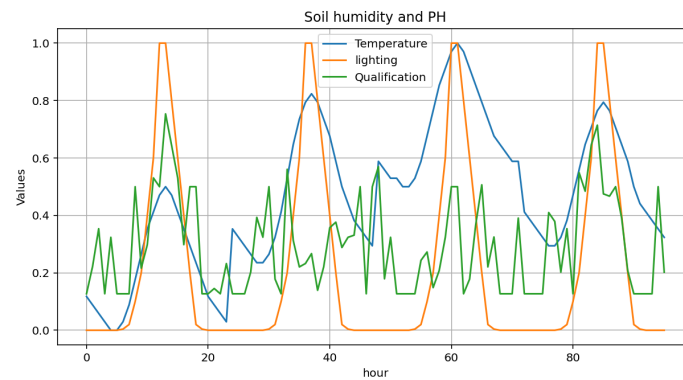


Fig. 4. Evaluation system based on a hierarchical fuzzy system.

- Soil moisture: 39.45
- Soil pH: 7.52

3.4 Considerations for Ideal Conditions

It is important to note that achieving these ideal conditions is very difficult unless additional measures are applied, such as fertilizers, irrigation systems, or greenhouse cultivation.

4 Conclusion

The nopal is a resilient plant that can adapt to less-than-ideal growing conditions. However, to maximize its growth and production, it is important to ensure that the conditions mentioned above are optimal. Monitoring and

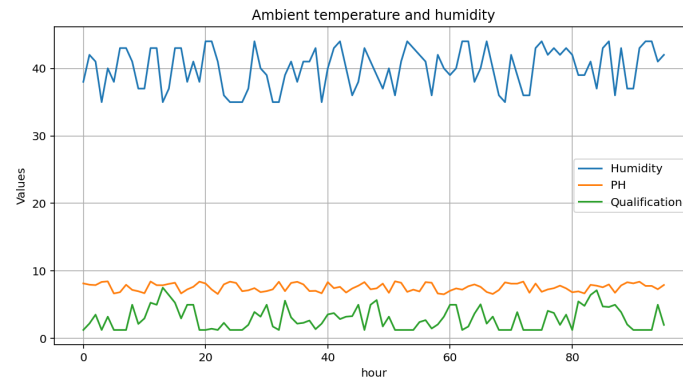


Fig. 5. Evaluation system based on a hierarchical fuzzy system.

adjusting these factors will significantly contribute to the success of nopal cultivation.

Exposure to air pollutants such as alcohols, methane, and hydrocarbons can be harmful to nopal crops, affecting their health, growth, and productivity. Implementing appropriate management strategies and continuous monitoring can help mitigate these effects and promote healthier, more productive cultivation.

Airborne particulate matter poses a significant risk to the health and productivity of nopal crops. Understanding and mitigating the impacts of this type of pollution is crucial to maintaining plantation health and the quality of derived products.

The presented module demonstrates the feasibility of implementing a crop evaluation system focused on nopal cultivation. The results obtained allow for application to extensive crops, as the system can be mounted on a robot and perform plant-by-plant inspections. The issue of environmental contamination remains a challenge, particularly regarding the use of more sophisticated sensors or for crops that require greater care. Finally, this system can be extended to other crops, considering modifications in the rules, as each type of plant has specific care requirements.

A hierarchical fuzzy system provides a robust framework for evaluating crop quality by systematically processing multiple attributes at different levels. This approach handles uncertainty and complexity effectively, making it suitable for use in portable devices, inclusive of drones or autonomous vehicles.

References

1. Abdullah, N., Durani, N. A. B., Shari, M. F. B., Siong, K. S., Hau, V. K. W., Siong, W. N., Ahmad, I. K. A.: Towards smart agriculture monitoring using fuzzy systems. *IEEE Access*, vol. 9, pp. 4097–4111 (2021) doi: 10.1109/ACCESS.2020.3041597

2. Baradaran, A. A., Tavazoei, M. S.: Fuzzy system design for automatic irrigation of agricultural fields. *Expert Systems with Applications*, vol. 210, pp. 118602 (2022) doi: <https://doi.org/10.1016/j.eswa.2022.118602>
3. Bin, L., Shahzad, M., Khan, H., Bashir, M. M., Ullah, A., Siddique, M.: Sustainable smart agriculture farming for cotton crop: A fuzzy logic rule based methodology. *Sustainability*, vol. 15, no. 18 (2023) doi: 10.3390/su151813874
4. Cavaliere, D., Senatore, S., Loia, V.: Crop health assessment through hierarchical fuzzy rule-based status maps. *Knowledge and Information Systems*, vol. 66, no. 11, pp. 7109–7136 (2024) doi: 10.1007/s10115-024-02180-w
5. Dhal, S., Wyatt, B. M., Mahanta, S., Bhattarai, N., Sharma, S., Rout, T., Saud, P., Acharya, B. S.: Internet of things (iot) in digital agriculture: An overview. *Agronomy Journal*, vol. 116, no. 3, pp. 1144–1163 (2024) doi: <https://doi.org/10.1002/agj2.21385>
6. Duguma, A. L., Bai, X.: How the internet of things technology improves agricultural efficiency. *Artificial Intelligence Review*, vol. 58, no. 2, pp. 63 (2024) doi: 10.1007/s10462-024-11046-0
7. Kumar, K., Singh, D., Singh, R.: Cactus Pear: Cultivation and Uses. CIAH/Tech./Pub. No. 73, ICAR-Central Institute for Arid Horticulture, Bikaner, Rajasthan, India (2018)
8. Kumar, K., Sharma, S., Pandey, P., Goyal, H. R.: Iot enabled crop detection system using soil analysis. In: 2022 7th International Conference on Communication and Electronics Systems (ICCES). pp. 323–327 (2022) doi: 10.1109/ICCES54183.2022.9835959
9. Puri, V., Chandramouli, M., Le, C. V., Hiep Hoa, T.: Internet of things and fuzzy logic based hybrid approach for the prediction of smart farming system. In: 2020 International Conference on Computer Science, Engineering and Applications (ICCSEA). pp. 1–5 (2020) doi: 10.1109/ICCSEA49143.2020.9132933
10. Rajak, P., Ganguly, A., Adhikary, S., Bhattacharya, S.: Internet of things and smart sensors in agriculture: Scopes and challenges. *Journal of Agriculture and Food Research*, vol. 14, pp. 100776 (2023) doi: <https://doi.org/10.1016/j.jafr.2023.100776>
11. Saini, A., Gill, N. S., Gulia, P., Tiwari, A. K., Maratha, P., Shah, M. A.: Smart crop disease monitoring system in iot using optimization enabled deep residual network. *Scientific Reports*, vol. 15, no. 1, pp. 1456 (2025) doi: 10.1038/s41598-025-85486-1
12. Torres-Trevino, L.: Evonorm: Easy and effective implementation of estimation of distribution algorithms. In: *Advances in Computer Science and Engineering*. pp. 75–83 (2006)

Traductor de texto audio para ayuda en la enseñanza de débiles visuales

Armando Martinez Rios, Adriana Sosa Rojas,
Ruben Naftanael Palomino Flores

Instituto Politécnico Nacional,
Escuela Superior de Ingeniería Mecánica y Eléctrica (ESIME) Zacatenco,
México

{armartinezr, asosaro, rubennaf03}@gmail.com

Resumen. Este trabajo muestra el desarrollo de una aplicación destinada a promover la educación inclusiva para personas con discapacidad visual en la Escuela Superior de Ingeniería Mecánica y Eléctrica (ESIME), unidad Zacatenco, del Instituto Politécnico Nacional (IPN). Diseñada como un traductor, convierte material didáctico como presentaciones, documentos en formato docx y PDF, en archivos de audio MP3. La aplicación desarrollada presenta una interfaz gráfica de usuario que ayuda a cargar sus documentos en formatos pptx, docx y pdf, tanto desde su computadora como desde Google Drive con el apoyo de un asistente virtual. Los documentos son procesados para convertir los caracteres en archivos de audio. Python es el lenguaje de programación principal debido a la gama de herramientas y bibliotecas disponibles y a su código abierto.

Palabras clave: Traductor, asistente, procesamiento de lenguaje natural.

Text-to-Audio Translator for Assistance in the Education of the Visually Impaired

Abstract. This work shows the development of an application aimed at promoting inclusive education for visually impaired people at the Escuela Superior de Ingeniería Mecánica y Eléctrica (ESIME), Zacatenco unit, of the Instituto Politécnico Nacional (IPN). Designed as a translator, it converts didactic material such as presentations, documents in docx and PDF format, into MP3 audio files. The developed application features a graphical user interface that helps upload documents in pptx, docx, and pdf formats, both from the user's computer and from Google Drive with the support of a virtual assistant. The documents are processed to convert the characters into audio files. Python is the main programming language due to the range of available tools and libraries and its open source nature.

Keywords: Translator, assistant, natural language processing.

1. Introducción

1.1. Debilidad visual

La debilidad visual es una deficiencia que afecta la capacidad de las personas en su agudeza y campo visual, motilidad ocular, distinción de colores o profundidad en los objetos a diferencia de una persona con ceguera que tiene nula visión de su entorno. Sin embargo aunque se tiene un poco de visión, la posibilidad de asistir a la escuela (desde su nivel básico hasta el nivel profesional), se ve afectada por muchos factores. La accesibilidad a los textos que suele ocuparse en las instituciones ya sea en formato digital o en formato físico, se convierte en un obstáculo más. Estevez, [4] define debilidad visual como cualquier restricción o carencia (resultado de una deficiencia) de la capacidad de realizar una actividad en la misma forma o grado que se considera normal para un ser humano. Se refiere a actividades complejas e integradas que se esperan de las personas o del cuerpo en conjunto, como pueden ser las representadas por tareas, aptitudes y conducta.

En México se promulgó el 29 de junio de 2021 la ley General de Educación donde se adicionaron diversas disposiciones para asegurar la educación inclusiva, tomando para ello los artículos 61 a 68. Agregado a esto, en el Diario oficial de la Federación (DOF) con fecha del 6 de enero del 2023, se menciona que para asegurar una educación inclusiva deben aplicarse métodos, materiales de apoyo y técnicas específicas propiciando la integración de personas con alguna discapacidad [14].

1.2. Políticas públicas en instituciones de educación superior

En el ámbito de las políticas públicas y el apoyo financiero para fomentar la inclusión en México, las Instituciones de Educación Superior (IES) deben alinearse con las directrices y normativas de las entidades decisorias, que son las responsables de asignar presupuestos y recursos adicionales. La modernización de los mecanismos gubernamentales ha motivado a las IES a impulsar la inclusión educativa mediante programas específicos como becas dirigidas a estudiantes de grupos vulnerables y fondos adicionales. Aunque algunas IES desarrollan programas de inclusión independientes de estos incentivos, muchas participan en concursos para obtener recursos adicionales, lo que refleja una dependencia parcial de estas ayudas. Según un informe de Vadillo y Casillas en 2017, solo 17 IES en México implementan programas de inclusión, distribuidas en varios estados del país. A nivel federal, una de las medidas gubernamentales para fomentar la inclusión educativa en la educación superior es el Programa para la Inclusión y la Equidad Educativa (PIEE), establecido en 2015 y aún en vigor. Desde su inicio, este programa ha asignado un promedio anual de 70,075,958 pesos, lo que representa el 0.06 % del presupuesto total destinado a la educación superior en 2018. Estos fondos se distribuyen anualmente entre las diversas IES públicas que compiten con proyectos de inclusión [6].

En la ESIME Zacatenco, no se cuenta con evidencia de algún programa de apoyo destinado a la elaboración de material didáctico o la adaptación de trámites administrativos dirigidos específicamente a personas con debilidad visual. Esta situación es notable debido a la naturaleza educativa de la institución, donde la falta de

estos recursos puede representar una barrera significativa para el acceso al conocimiento por parte de la población estudiantil que presenta esta discapacidad. Además, la ausencia de herramientas adecuadas en los trámites administrativos puede tener un impacto negativo en el proceso académico. Aunque la población con discapacidad visual en el instituto pueda considerarse como una minoría, tanto a nivel nacional como internacional existen normativas que promueven los derechos de estas personas a la educación inclusiva.

En este contexto, surge la necesidad de desarrollar un programa informático que apoye a los docentes en la traducción de sus textos a formato de audio, facilitando así la inclusión de estudiantes con debilidad visual en el ámbito educativo de la ESIME Zacatenco. La carencia de herramientas específicas y adaptadas a las necesidades de esta población en la institución evidencia la urgencia de abordar este problema y contribuir a la mejora del acceso a la educación para todos los estudiantes, independientemente de sus habilidades visuales. Por lo que surge la siguiente pregunta **¿Cómo desarrollar un asistente para que por medio de reconocimiento óptico de caracteres que permita la traducción de texto a audio en estudiantes con debilidad visual?**

1.3. Objetivo general

Desarrollar un prototipo para traducción de información escrita a audio basada en un asistente virtual que ayude a estudiantes de nivel superior con deficiencia visual

2. Estado del arte

El trabajo *"Soluciones digitales de inclusión educativa para personas con discapacidad visual"*, describe la discapacidad visual, abarcando sus tipos, características, causas y prevalencia global. El estudio también examina los obstáculos que enfrentan las personas con discapacidad visual para acceder a la educación superior y explora cómo el desarrollo tecnológico, mediante el uso de herramientas digitales, puede ofrecer soluciones [7]. Mientras que en *"Implementación de un asistente virtual para los estudiantes de pregrado de una universidad peruana"* se muestra la implementación de un asistente virtual destinado a los estudiantes mencionados. Se plantearon tres objetivos específicos: analizar el tiempo de respuesta del asistente, evaluar su eficiencia, y comprobar la percepción de los estudiantes sobre su uso. La evaluación del asistente se realizó en un enfoque cuantitativo con un diseño no experimental y de nivel descriptivo. La población objetivo consistió en 3,228 estudiantes, de los cuales se seleccionó una muestra de 78 estudiantes de ingeniería industrial [11].

El dispositivo vestible lector de textos impresos para niños con discapacidad visual se concibe como un recurso de apoyo en el contexto educativo. Aprovecha la visión artificial, la tecnología de reconocimiento óptico de caracteres (OCR) y la síntesis de voz (TTS) para ofrecer acceso a la información escrita, transformándola en información hablada [8]. Mientras que en España, un equipo desarrolló una plataforma para mejorar la calidad educativa, en la cual implementaron la conversión de texto a

audio (TTS). Su objetivo era crear un entorno digital que ofreciera diversos servicios, permitiendo capturar la información transmitida por el profesor en forma de voz para luego transformarla en datos digitales que pudieran ser procesados, modificados o compartidos con los distintos dispositivos utilizados por los estudiantes. De esta manera, los estudiantes podrían acceder a una representación tanto escrita como auditiva de las clases impartidas por el profesor [15].

El artículo *"Asistente Virtual Académico Utilizando Tecnologías Cognitivas De Procesamiento De Lenguaje Natural"* proporciona un análisis del estado actual del uso de tecnologías de Procesamiento de Lenguaje Natural (PLN) en el desarrollo de asistentes virtuales académicos. Los autores detallan el uso de Dialogflow, una plataforma de desarrollo de agentes conversacionales de Google, como base para la creación del asistente virtual. Destacan cómo esta herramienta permite la interacción fluida entre el asistente y los estudiantes, capacitándolo para responder a una amplia gama de consultas [12].

En el Estado de Puebla, México, opera un grupo académico conocido como ITTEH-CA-09, que ha estado trabajando durante los últimos tres años en la investigación y desarrollo de software destinado a fomentar la inclusión de personas con discapacidad visual y auditiva. Han desarrollado una aplicación que utiliza tecnologías accesibles, como dispositivos móviles, para la captura y procesamiento de imágenes, lo que facilita el acceso de personas con discapacidad visual a entornos educativos y les ayuda a comprender textos con nomenclatura matemática [9].

El artículo *"Asistente Virtual Académico Utilizando Tecnologías Cognitivas De Procesamiento De Lenguaje Natural"* proporciona un análisis exhaustivo del estado actual del uso de tecnologías de Procesamiento de Lenguaje Natural (PLN) en el desarrollo de asistentes virtuales académicos. Los autores detallan el uso de Dialogflow, una plataforma de desarrollo de agentes conversacionales de Google, como base para la creación del asistente virtual. Asimismo, el artículo resalta la integración del asistente virtual con Fulfillment, una herramienta que posibilita la generación dinámica de respuestas y la ejecución de acciones contextuales [12].

En síntesis la comunidad científica ha estado interesada en desarrollar herramientas tecnológicas que apoyen en la educación y en la vida misma de comunidades vulnerables, sin embargo no han permeado lo suficiente y en la ESIME no se tiene ningún tipo de este apoyo.

3. Marco teórico

3.1. Débil visual

La debilidad visual surge cuando una condición ocular afecta al sistema visual y sus funciones relacionadas con la visión. Esta afección puede ser el resultado de diversas enfermedades oculares, como el debilitamiento de los vasos sanguíneos que alimentan la retina, el daño al nervio óptico debido al aumento de la presión ocular, lesiones oculares, o incluso puede ser una condición congénita. A nivel global, las principales causas de debilidad visual y ceguera son los errores de refracción y las cataratas. Los errores de refracción implican problemas con la forma del ojo que

causan una refracción incorrecta de la luz, resultando en visión borrosa. Por otro lado, las cataratas son opacidades en el cristalino del ojo que obstruyen el paso de la luz a la retina [13]. Se estima que al menos 2,200 millones de personas en el mundo presentan algún tipo de deterioro visual, de las cuales al menos 1,000 millones podrían haber prevenido o tratado su discapacidad visual. Sin embargo, solo el 36 % de las personas con deterioro visual por errores de refracción y el 17 % de aquellas con discapacidad visual a causa de cataratas han tenido acceso a una intervención adecuada. La discapacidad visual representa una carga económica considerable a nivel global, con un costo anual en términos de productividad que se estima en 411,000 millones de dólares estadounidenses [10].

3.2. Educación inclusiva

La educación inclusiva es un enfoque pedagógico que busca garantizar que todos los estudiantes, sin importar sus características individuales, tengan acceso y participación en un entorno de aprendizaje que promueva la igualdad, la diversidad y la tolerancia. Este enfoque implica una revisión profunda de las políticas, prácticas y cultura de las instituciones educativas, y necesita la participación activa de toda la comunidad escolar. En la educación inclusiva, se busca eliminar las barreras que puedan restringir el acceso y la participación de los estudiantes, ya sea por discapacidad, género, origen étnico, orientación sexual u otras diferencias. Se fomenta un entorno en el que todos los alumnos se sientan valorados, respetados y apoyados en su proceso de aprendizaje, así como el desarrollo de habilidades sociales y emocionales para convivir en la diversidad. La educación inclusiva no solo beneficia a los estudiantes con necesidades particulares, sino que también enriquece la experiencia educativa de todos al promover la comprensión mutua, el respeto por las diferencias y la colaboración entre compañeros. Es un proceso continuo que requiere un compromiso constante hacia la equidad y la justicia en la educación [13].

3.3. Aplicaciones de software y herramientas tecnológicas en el ámbito educativo

El empleo de herramientas tecnológicas en el ámbito educativo abarca desde la creación de plataformas web para la formación de estudiantes, hasta la utilización de la inteligencia artificial para adaptar el proceso de aprendizaje, enriquecer la experiencia educativa, automatizar tareas administrativas y analizar datos con el fin de mejorar la experiencia estudiantil en la escuela [2]. Soluciones digitales de inclusión, que aprovechan el avance de la tecnología para mejorar la experiencia educativa de las personas con discapacidad visual. Estas herramientas facilitan una mayor interacción entre estudiantes y profesores, permitiendo diseñar estrategias de trabajo adaptadas a diferentes estilos de aprendizaje y modelos educativos inclusivos [7].

En el caso de hardware, se han desarrollado dispositivos como los de tecnología vestible para asistir el desplazamiento de personas con discapacidad visual, mejorando así su orientación y seguridad en entornos interiores y en algunos casos en exteriores [5].

3.4. Lenguaje Natural y su procesamiento

El Procesamiento del Lenguaje Natural (PLN) es una área de la inteligencia artificial que se enfoca en facilitar la comunicación entre las computadoras y el lenguaje humano. Su objetivo principal es dotar a las máquinas de la capacidad de entender, interpretar y generar lenguaje humano de forma natural. El PLN tiene diversas aplicaciones prácticas, como la traducción automática, la creación de resúmenes, la clasificación de textos y la respuesta automática a preguntas. Para alcanzar estos fines, el PLN se basa en técnicas de disciplinas como la lingüística computacional, el aprendizaje automático y la minería de datos, lo que permite procesar grandes volúmenes de datos en lenguaje natural y extraer información relevante [3]. La Generación Natural del Lenguaje (NLG) es el proceso mediante el cual los sistemas informáticos producen texto de manera automática, con una estructura y contenido que se asemejan al lenguaje humano. Este proceso implica varias tareas, como determinar el contenido del texto, estructurarlo, seleccionar palabras y frases, agregar información relevante, generar expresiones de referencia, asegurar la gramaticalidad y presentar el texto de manera coherente [1].

3.5. Población

La Escuela La Superior de Ingeniería Mecánica y Eléctrica (ESIME) en su unidad Zacatenco, es una de las escuelas fundadoras del Instituto Politécnico Nacional (IPN), con una matrícula de 12,500 estudiantes inscritos para el periodo 2024/2 en dos turnos; este número se divide en cinco carreras: Ingeniería en Comunicaciones y Electrónica (ICE), con casi la mitad de la matrícula; le siguen ingeniería eléctrica (IE) e ingeniería en control y automatización (ICA) con otro 40 % y el otro 10 % dividido en ingeniería en sistemas automotrices (ISISA) e Ingeniería Fotónica (IF) de nueva creación. Esta cantidad de estudiantes la ubica entre las tres escuelas más grandes, en cuanto a población de todo el Instituto Politécnico Nacional (IPN). Esta matrícula se fue incrementando de forma paulatina de 2004 a 2015. En lo que respecta a docentes, entre las cuatro carreras se tiene casi un millar de ellos, de los cuales alrededor de 600 dan servicio a la carrera de ICE y el resto se divide en las otras cuatro carreras. Hasta esta fecha no se tiene un dato exacto oficial de la población con problemas visuales pero si se tiene evidencia no oficial a través de los compañeros docentes.

4. Propuesta de diseño

El sistema propuesto se compone de cinco etapas como son el asistente virtual, la interfaz gráfica de usuario, la conversión, la integración con Google y el empaquetado. Las etapas se muestran en la figura 1. Estas etapas no necesariamente ocurren una después de otra sino que se dividieron así para una mejor explicación.

4.1. Asistente virtual

Esta etapa fué diseñada para interactuar con los usuarios y guiarlos a través del proceso de la selección de archivos para la conversión en audio. Cuando el



Fig. 1. Diagrama a bloques de la propuesta.

usuario inicia el asistente, se le da la bienvenida con un mensaje introductorio que explica las funcionalidades disponibles y las opciones para comenzar. Este asistente está programado para responder a los comandos de voz del usuario utilizando el reconocimiento de voz para interpretar las solicitudes y preguntas del usuario.

Una de las primeras decisiones que el asistente guía al usuario a tomar es la selección de la ubicación del archivo que desea convertir. Esta ubicación puede ser de forma local en el dispositivo del usuario o en alguna cuenta de Google Drive. El asistente recoge la respuesta y procede según la opción seleccionada. Si se elige la opción de la computadora local, el asistente muestra las rutas disponibles y permite al usuario seleccionarla; Luego, muestra los archivos disponibles en esa ruta y permite al usuario elegir el archivo que desea convertir en audio.

Por otro lado, si el usuario elige la opción de Google Drive, el asistente solicita la autenticación y el acceso a la cuenta de Google Drive del usuario. Una vez autenticado, el usuario puede proporcionar un enlace al archivo que desea convertir, y el asistente descarga automáticamente el archivo para su conversión.

4.2. Interfaz gráfica de usuario

La interfaz de la aplicación de conversión de texto a audio se ha diseñado con varias pantallas y elementos interactivos para proporcionar al usuario una experiencia fluida y eficiente. Al iniciar la aplicación, la pantalla de portada se presenta como el punto de entrada, ofreciendo dos opciones principales: "Conversión Manual" o "Conversión por Asistente Virtual". Cada opción se presenta mediante botones que el usuario puede seleccionar con un clic. Si el usuario elige la opción "Conversión Manual", se despliega la pantalla correspondiente (InitialScreen), donde puede seleccionar un archivo de texto de su sistema local. Esta pantalla facilita la selección del archivo mediante un botón de búsqueda, habilita una ventana de diálogo para que el usuario elija el archivo deseado.

En caso de elegir la opción de "Conversión por Asistente Virtual", la pantalla de Asistente Virtual se inicia. Aquí, el usuario es guiado a través del proceso de conversión por el asistente virtual que utiliza reconocimiento de voz. El asistente le permite seleccionar la ubicación del archivo que desea convertir, ya sea en su computadora o en Google Drive, y luego inicia la conversión a audio.

Durante el proceso de conversión, la pantalla de Conversión (PowerPointToAudioConverter) proporciona una vista previa de su contenido y un botón para iniciar la conversión a formato de audio. Una vez que la conversión se completa, se ofrece al usuario la opción de reproducir el archivo de audio resultante para su revisión. Además, la aplicación incluye una Ventana de Ruta (VentanaRuta) que permite al usuario ingresar rutas de archivos locales y enlaces de Google Drive para facilitar la búsqueda y selección de archivos en futuras sesiones.

4.3. Conversión

Esta etapa consiste en la extracción del contenido de archivos de texto o imagen mediante un script, donde se se abren los archivos en formatos como pptx, PDF o docx, y se extrae información de sus diapositivas o páginas según corresponda. Luego, se lleva a cabo el reconocimiento óptico de caracteres (OCR) en las imágenes y, finalmente, se procedía a la lectura en voz alta del contenido de cada diapositiva o página.

El programa correspondiente escrito para el lenguaje python V3.x; consta de dos clases, la clase principal que efectúa el reconocimiento del texto y lo convierte a audio y la clase que muestra el estado de la conversión mediante una interfaz. La clase `AudioConversionThread`, hereda de `QThread` y gestiona las tareas en segundo plano. Su método `init` inicializa la clase con la información necesaria, mientras que el método `run` es donde se lleva a cabo el trabajo principal. Este trabajo implica acciones como extraer texto de archivos PDF, convertir documentos DOCX a PDF, o extraer información de presentaciones PowerPoint, dependiendo de la extensión del archivo.

Por otro lado, la clase `PowerPointToAudioConverter` hereda de `QWidget` y se encarga de interactuar con la interfaz gráfica de usuario. Su método `init` inicializa la clase y establece las variables necesarias, mientras que el método `init_ui` configura la interfaz gráfica.

Estas clases están conectadas de manera que cuando se inicia la conversión de audio desde la interfaz gráfica, se crea una instancia de `AudioConversionThread` en un hilo separado. La señal `finished` de `AudioConversionThread` está conectada a la función `conversion_completed` de `PowerPointToAudioConverter`, lo que permite actualizar la interfaz gráfica de usuario según el estado de la conversión.

4.4. Integración con Google

Tras el desarrollo inicial del asistente virtual, se integró la funcionalidad para gestionar archivos almacenados directamente en la nube, específicamente en Google Drive con el fin de brindar un mayor apoyo al usuario en la obtención de los archivos que desea convertir.

La integración de Google Drive como una opción de almacenamiento proporciona a los usuarios un acceso a sus archivos desde cualquier dispositivo con conexión a internet. Ahora, el asistente virtual permite al usuario seleccionar archivos tanto de su sistema local como de su cuenta de Google Drive. Esta mejora busca que la experiencia del usuario al hacer que el proceso de conversión de texto a audio sea más rápida.

4.5. Empaquetado

Auto-Py-to-Exe es una herramienta que permite convertir scripts de Python en ejecutables independientes, incluyendo todas las bibliotecas y archivos necesarios para su correcto funcionamiento. Este proceso de empaquetado es necesario para garantizar que la aplicación sea fácilmente distribuible y ejecutable en sistemas que no tengan instalado Python ni las bibliotecas utilizadas. La creación del ejecutable con Auto-Py-to-Exe generalmente implica seleccionar el script principal de la aplicación, así como especificar cualquier archivo adicional que sea necesario incluir, ya sea



Fig. 2. Pantalla de bienvenida al iniciar el asistente.

imágenes, archivos de configuración o bases de datos. Después, AutoPy-to-Exe se encarga de analizar las dependencias del script y empaquetar todo en un único archivo ejecutable que puede ser distribuido y ejecutado en sistemas Windows sin necesidad de tener instalado Python. Ya generado el archivo ejecutable, la aplicación está lista para ser distribuida a los usuarios finales, lo que facilita su uso en diferentes entornos sin requerir configuraciones adicionales.

5. Resultados

La versión final de la aplicación desarrollada con PyQt5 cuenta con las dos opciones de conversión ya comentadas, manual y asistida; llamada asistente virtual, como se muestra en la figura 2.

La conversión manual permite a los usuarios buscar o arrastrar el archivo deseado directamente en la aplicación, ofreciendo una interacción intuitiva y eficiente. La opción asistida, por otro lado, permite cargar el archivo mediante comandos de voz, ya sea desde la computadora o desde un servicio de almacenamiento en la nube, (Google Drive). La ventana de salida se muestra en la figura 3.

Para simplificar la búsqueda de archivos, se pueden añadir rutas específicas de la computadora y estas se van a ir guardando por si se requieren nuevamente o modificar el enlace de Google Drive para seleccionar otro archivo. Además, se integran algoritmos para el Reconocimiento Óptico de Caracteres y la conversión de texto a voz; lo que permite extraer y procesar datos de archivos en formato PPTX, PDF y DOCX, incluyendo tanto texto como imágenes.

Utilizando el motor de Tesseract y el lenguaje de programación Python, el OCR transforma las imágenes de las presentaciones o de los archivos en un formato legible por máquina, que posteriormente se convierte en audio MP3. Este formato de salida es ampliamente compatible y facilita la reproducción en diversos dispositivos compatibles con el mismo.

Cuando se realiza una conversión automática dentro de la aplicación, el usuario tiene que seleccionar el archivo, ya sea PDF, PPTX o DOCX y no realizar ningún



Fig. 3. Pantalla de salida de conversión en proceso.

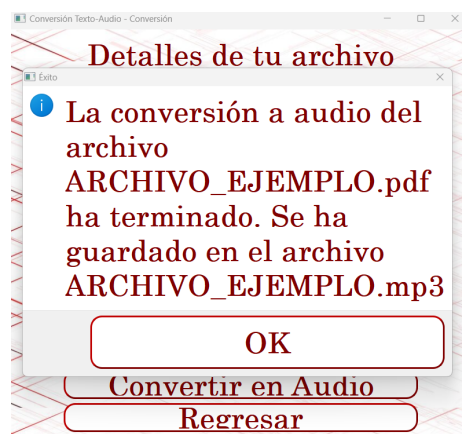


Fig. 4. Pantalla de fin del proceso.

proceso técnico subyacente. Básicamente, el usuario solo necesita saber la ubicación de su archivo en el dispositivo o en su Google Drive.

Los archivos PDF, PPTX y DOCX son formatos ampliamente utilizados en el ámbito educativo, lo que garantiza que la aplicación sea compatible con los tipos de documentos más comunes empleados por los docentes y estudiantes. El archivo de audio resultante se guarda automáticamente en la misma ubicación que el archivo ejecutable de la aplicación, lo que facilita la tarea de búsqueda para el usuario. Además, la interfaz de la aplicación permite reproducir directamente el archivo una vez que ha sido creado. Cuando termina el proceso se informa al usuario con la ventana mostrada en la figura 4.

La aplicación es completamente portátil y puede utilizarse en cualquier dispositivo con sistema operativo Windows 10 o superior. Sin embargo, es importante destacar que su funcionamiento requiere una conexión a internet, lo que permite acceder a los

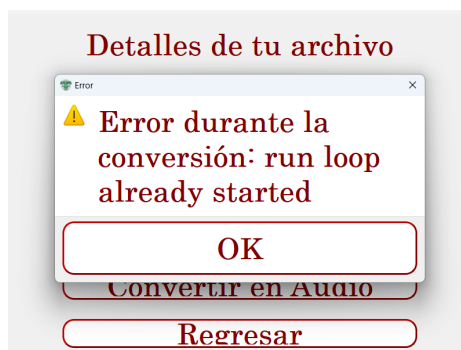


Fig. 5. Pantalla de salida de conversión en proceso.

servicios en la nube y a las actualizaciones necesarias para la conversión de archivos y el procesamiento de datos. Esta combinación de portabilidad y conectividad asegura que los usuarios puedan aprovechar al máximo las capacidades de la aplicación en cualquier momento y lugar, siempre que dispongan de acceso a la red.

Cabe señalar que el asistente está programado para manejar errores durante el reconocimiento de voz, como el que se muestra en la figura 5. Si el usuario menciona un número de ruta o documento fuera de rango, como "ruta siete" por ejemplo, cuando solo hay cinco rutas disponibles, el asistente informará que la selección está fuera de rango y pedirá una nueva selección. También, si no se reconoce la selección de la ruta o el documento, el asistente pedirá nuevamente al usuario que diga la ruta o el documento seguido de su número. El asistente también maneja errores en la detección de audio, ya sea por que el usuario esté hablando muy bajo o porque hubo algún ruido externo que impidió que se escuchara correctamente lo que el usuario dijo, si el asistente no pudo detectar correctamente el audio, dará el siguiente mensaje: "No se detectaron voces claras. Por favor, hable más fuerte y vuelva a intentarlo".

El asistente tiene un límite de intentos para seleccionar rutas y documentos. Si se excede este límite, el asistente despliega un mensaje que indica "Se ha excedido el número de intentos. Por favor, inicia de nuevo el asistente", y el usuario deberá presionar nuevamente el botón de Asistente Virtual para reiniciar el proceso, esto en caso de no poder seleccionar un documento o ruta, o no se haga bien el reconocimiento de voz el programa no siga preguntando de manera continua y detener el proceso. El asistente también incluye una opción para cancelar el proceso en cualquier momento. Si el usuario dice la palabra "cancelar", el asistente dará el mensaje: "Asistente cancelado. Por favor, inicia de nuevo si deseas utilizarlo", y el usuario deberá presionar nuevamente el botón de Asistente Virtual para reiniciar el asistente.

6. Conclusiones

Este proyecto ha alcanzado con éxito el desarrollo de una herramienta de software en forma de un asistente virtual, para convertir material didáctico en formato de audio MP3 en el sistema operativo Windows 10 o superior; cumpliendo la meta al integrar algoritmos como el Reconocimiento Óptico de Caracteres y la conversión de texto

a voz, facilitando así la interpretación eficiente de archivos multimedia. Además, el desarrollo de un asistente virtual es capaz de reconocer voz y ejecutar comandos específicos relacionados con la conversión de material educativo ha enriquecido significativamente la funcionalidad de la herramienta. La implementación exitosa de tecnologías como Tesseract OCR y pyttsx3 ha permitido crear archivos de audio MP3 a partir de contenido textual y gráfico.

Esta aplicación ofrece una solución para la conversión de contenido educativo visual en narraciones auditivas, apoyando a docentes y estudiantes en el proceso de enseñanza y aprendizaje inclusivo. La implementación de múltiples rutas de búsqueda y la compatibilidad con archivos almacenados en la nube mejoran la flexibilidad y accesibilidad del software, al tener disponibles los archivos de conversión en todo momento.

Se cumplió el objetivo planteado al lograr desarrollar una plataforma para conversión de texto a audio en el lenguaje de programación Python y compatible con el sistema operativo Windows 10 o superior; alcanzando el registro de derechos de autor no. 03-2024-122013344100-01 por el Instituto Nacional de Derechos de autor.

Referencias

1. Bernardos-Galindo, M. d. S.: ¿Qué es la generación de lenguaje natural? Una visión general sobre el proceso de generación. *Inteligencia Artificial, revista Iberoamericana de inteligencia artificial*, vol. 11, pp. 105–128 (2007) doi: 10.4114/IA.V11I34.911.
2. Berrones, L., Salgado, S.: La aplicación de la inteligencia artificial para mejorar la enseñanza y el aprendizaje en el ámbito educativo. *Esprint Investigación*, vol. 2, pp. 52–60 (2023) doi: 10.61347/ei.v2i1.52.
3. Durán-Chinchilla, C. M., García-Quintero, C. L., Rosado-Gómez, A. A.: Aplicación del procesamiento del lenguaje natural como técnica de análisis en la producción textual, caso estudiantes de Ingeniería de Sistemas UFPSO. *Revista Boletín Redipe*, vol. 11, no. 1, pp. 489–497 (2022) doi: 10.36260/rbr.v11i1.1656.
4. Estévez-Alarcón, J. A., Ortiz-Vázquez, P. E., Ticante-Calderón, C. E.: Discapacidad visual como causa de depresión en adolescentes. *Ciencias de la Salud en la Ibero Puebla*, (2017)
5. Fariña-Molina, N. R., Ibarrola-Chamorro, C. R., Ramirez-Aquino, A. T.: Desarrollo tecnológico de un brazalete sensorizado (CiegNest) para apoyo a personas con discapacidad visual implementando software y hardware libre. In: *Anais do XIX Congresso Latino-Americano de Software Livre e Tecnologias Abertas*, pp. 149–152 (2022) doi: 10.5753/latinoware.2022.228046.
6. González-García, J. A., Zúñiga-Llamas, A., Arce-Casas, P. O.: Un panorama sobre la cobertura educativa a nivel superior en México para personas con discapacidad. *IE Revista de Investigación Educativa de la REDIECH*, vol. 12, pp. 1–17 (2021) doi: 10.33010/ie_rie_rediech.v12i0.1171.
7. Hernández-Bermejo, M. A., Leon-Pacheco, D., Minchala-González, J., Lema-Lema, E.: Soluciones digitales de inclusión educativa para personas con discapacidad visual. *Killkana Técnica*, vol. 7, pp. 19–32 (2024) doi: 10.26871/killkanatecnica.v7i2.1067.
8. Mejía-Campos, R. G.: Dispositivo Wereable lector de textos impresos para niños con discapacidad visual como herramienta de apoyo en el proceso de enseñanza y aprendizaje en el ámbito escolar usando técnicas de visión artificial, OCR y TTS. Tesis de Ingeniería, Benemérita Universidad Autónoma de Puebla, (2022)

9. Olgún-Gil, L. E., Vázquez-Zayas, E., Vázquez-Guzmán, F., García-Ortega, I., Rodríguez-Ramírez, F.: Aplicación móvil que traduce ecuaciones matemáticas a voz mediante OCR. *Pistas Educativas*, vol. 42, pp. 74–90 (2020)
10. OMS: Ceguera y discapacidad visual, (Feb 2024)
11. Ortiz-Elías, A. N., Dávila-Morán, R. C.: Implementación de un asistente virtual para los estudiantes de pregrado de una universidad peruana. *Revista Conrado*, vol. 19, pp. 121–128 (2023)
12. Torres-Echeverri, M. M., Manjarrés-Betancur, R.: Asistente virtual académico utilizando tecnologías cognitivas de procesamiento de lenguaje natural. *Revista Politécnica*, vol. 16, no. 31, pp. 85–96 (2020) doi: 10.33571/rpolitec.v16n31a7.
13. Vallejos-Garcías, V., Castro-Durán, L.: Educación inclusiva. *Revista Internacional de Humanidades*, vol. 12, pp. 1–11 (2023) doi: 10.37467/revhuman.v12.4726.
14. Villa-Guardiola, V. J., Lorenzo-Loredo, A. G., Sayago Gonzalez, I. N.: Educación superior incluyente para personas con discapacidad visual en la Universidad Autónoma de Guerrero, México. *Saber, Ciencia y Libertad*, vol. 17, pp. 474–493 (2022) doi: 10.18041/2382-3240/saber.2022v17n2.9332.
15. Villarrubia-González, G., de-Paz-Santana, J. F., Moreno-García, M. N., et al.: E-SPEECH: plataforma de apoyo a la comprensión oral mediante tecnologías disruptivas, (2022)

Identificación de patrones lingüísticos, temáticos, emocionales e intención en los corridos y la música regional mexicana

José Eder Martínez-Martínez ¹, Yessenia Díaz-Álvarez ¹,
Tlálloc Humberto Vergara-Morales ¹, Víctor Manuel Millán-Aguilar ²,
Gabriel González-Serna ¹, Noé Alejandro Castro-Sánchez ¹

¹ Tecnológico Nacional de México, campus CENIDET,
México

² Tecnológico Nacional de México, campus Zacatepec,
México

{m24ce005, d24ce109, m23ce088, gabriel.gs, noe.cs}
@cenidet.tecnm.mx, L20091134@zacatepec.tecmm.mx

Resumen. La música regional mexicana, especialmente los corridos, ha experimentado una evolución significativa en su contenido lírico, reflejando influencias culturales, sociales y digitales. Este estudio realiza un análisis computacional integral de los corridos, combinando modelado de tópicos, análisis de sentimientos y clasificación de intención. A través del Reconocimiento de Entidades Nombradas (NER), se identifican términos recurrentes y referencias geográficas que evidencian la influencia de la migración y la globalización. Además, mediante BERTopic, se extraen los temas predominantes y se analizan las tendencias temáticas. Para evaluar la carga emocional de las letras, se emplea una variante del modelo RoBERTa-base-bne, permitiendo la detección automática de sentimientos en los textos, analizando los cambios lingüísticos y patrones recurrentes. Finalmente, empleando BART-large-MNLI, se clasifican las canciones según si glorifican, desprestigian, o simplemente describen actividades delictivas. Los resultados proporcionan una visión profunda de los mensajes transmitidos a través de los corridos y sus implicaciones culturales y sociales.

Palabras clave: Procesamiento de lenguaje natural, modelado de tópicos, corridos mexicanos, análisis de sentimientos y emociones, clasificación de intención.

Identifying Linguistic, Thematic, Emotional and Intent Patterns in Corridos and Mexican Regional Music

Abstract. Mexican Regional Music, particularly “corridos”, have experimented a major evolution within its lyrical content, reflecting new cultural, social and digital influences. This study conducts a comprehensive computational analysis of “corridos”, combining topic modeling, sentiment analysis, and intent classification. Through Named Identity Recognition (NER), recurring terms and

geographical references are identified, evidencing the influence of migration and globalization. Additionally, through BERTopic, the predominant themes are extracted and the thematic trends are analyzed. To evaluate the emotional load of the lyrics, a variant of the RoBERTa-base-bne model is used, allowing the automatic detection of sentiments in the texts, analyzing linguistic changes and recurring patterns. Finally, using BART-large-MNLI, the songs are classified according to whether they glorify, vilify or simply describe criminal activities. The results provide an in-depth view of the messages conveyed through “corridos” and their cultural and social implications.

Keywords: Natural language processing, topic modeling, mexican corridos, sentiment and emotion analysis, intent classification.

1. Introducción

Los corridos mexicanos han sido un medio de narración que ha evolucionado junto con los cambios sociales y culturales, abordando desde historias de héroes revolucionarios hasta figuras del crimen organizado. Su contenido refleja influencias de la globalización y la era digital, transformando su narrativa para conectar con nuevas audiencias.

Este estudio emplea técnicas de Procesamiento de Lenguaje Natural (PLN) para analizar la estructura y evolución temática de los corridos y la música regional mexicana. Mediante BERTopic, se identifican los temas predominantes, mientras que el Reconocimiento de Entidades Nombradas (NER) permite detectar patrones lingüísticos clave, como nombres propios, lugares y términos recurrentes que reflejan las principales narrativas del género. Además, el análisis de sentimientos con RoBERTa-base-bne permite evaluar la carga emocional de las letras, proporcionando una perspectiva más profunda sobre su impacto.

Por último, la clasificación de intención con BART-large-MNLI ayuda a distinguir entre canciones que glorifican el crimen y aquellas que simplemente lo narran. A través de esta combinación de técnicas, se busca comprender mejor el papel de los corridos en la música y la sociedad contemporánea.

2. Trabajos relacionados

El Reconocimiento de Entidades Nombradas (NER) y el análisis de frecuencia de términos han sido utilizados para estudiar patrones temáticos en textos musicales. Investigaciones previas han aplicado estas técnicas en distintos enfoques [1, 2], pero su uso en los corridos y la música regional mexicana sigue siendo limitado. El análisis de frecuencia léxica ha demostrado ser útil para identificar tendencias en la música [3], permitiendo comprender cambios en el contenido lírico. Este estudio combina NER y análisis de términos para explorar la evolución temática de los corridos, aportando una visión cuantitativa sobre sus principales narrativas.

Existen diversos estudios sobre el impacto de los géneros musicales en la sociedad y la detección de temáticas en letras de canciones. Investigaciones previas han aplicado técnicas de modelado de tópicos para analizar narrativas en la música popular y su

relación con estereotipos de género y violencia. Sin embargo, pocos estudios se han centrado específicamente en los corridos. El modelado de temas agrupa documentos para resumirlos o clasificarlos. Al aplicarse a las letras de las canciones, proporciona un método eficaz para identificar temas recurrentes [4, 5]. Algoritmos de modelado de temas como BERTopic se han utilizado previamente en la investigación de género y ciencias sociales, y Wickham [6] utilizó este algoritmo para estudiar las expectativas de género en las redes sociales y su impacto en la ideación suicida.

Algunos estudios han explorado cómo estas canciones influyen en la percepción de la violencia y el crimen en la sociedad, mientras que otros han examinado su papel como una forma de protesta o resistencia cultural [7, 8]. Además, en el campo de la lingüística y el análisis de contenido, el Procesamiento de Lenguaje Natural (PLN) ha permitido analizar los temas y el sentimiento de las letras de manera automática, proporcionando una perspectiva sobre las emociones predominantes y la narrativa detrás de estos temas [9].

En cuanto a la clasificación temática e inferencia de intención en los corridos, BART-large-MNLI ha sido utilizado para tareas de categorización de texto sin necesidad de entrenamiento adicional. Biswas [10] desarrolló una API utilizando Flask que integra este modelo, facilitando su aplicación en entornos de producción. Un modelo derivado, BART-Large-MNLI-Yahoo-Answers, ha sido ajustado específicamente para la categorización de temas en el conjunto de datos de Yahoo Answers, demostrando una alta precisión en la clasificación de textos, incluso en categorías no vistas durante el entrenamiento [11]. Además, se han discutido implementaciones prácticas de BART-large-MNLI en foros como Stack Overflow, abordando su despliegue en producción y su integración en interfaces de usuario [12]. Estas discusiones reflejan el creciente interés por el uso de modelos de inferencia de lenguaje natural en aplicaciones prácticas.

3. Metodología

En esta sección se detallan los procedimientos utilizados para recolectar, procesar y analizar los datos. Que se pueden encontrar en el siguiente enlace <https://github.com/JoseEderMartinezMartinez/AnalisisComputacionalCorridosyMusicaRegionalMexicana>

a. Recopilación de datos

Para la compilación del corpus, se implementó un sistema de web scraping con Scrapy, obteniendo letras de canciones de los géneros corrido y música regional mexicana desde la plataforma letras.com.

El proceso de scraping se estructuró en varias etapas:

1. Obtención de artistas por género: Se utilizó un archivo maestro en formato JSON con una lista de artistas de los géneros objetivo.
2. Extracción de enlaces de canciones: Se recorrieron las páginas de cada artista para recolectar los enlaces de sus canciones.

3. Descarga de letras: Se accedió a cada enlace de canción y se extrajo el contenido textual de las letras.

Para evitar restricciones en la extracción masiva, se integró un sistema de proxies dinámicos, obtenidos y validados previamente desde free-proxy-list.net. Adicionalmente, se implementaron filtros para eliminar letras duplicadas y registros incompletos.

Este método permitió la construcción de un corpus diverso y representativo de los géneros analizados, la identificación de patrones en las letras de corridos y música regional mexicana.

b. Limpieza de datos

Una vez recopiladas las letras de canciones, se llevó a cabo un proceso de limpieza para mejorar la calidad del corpus y asegurar que los datos fueran adecuados para el análisis de tópicos. Este preprocesamiento incluyó varias etapas, con el objetivo de eliminar ruido y normalizar los textos.

Eliminación de Datos No Relevantes. Se aplicaron filtros para descartar aquellas canciones que presentaban características no deseadas:

- Letras vacías o compuestas únicamente por espacios en blanco.
- Textos que contenían exclusivamente caracteres especiales o números sin contenido léxico significativo.
- Canciones con errores en la estructura del archivo JSON o con registros incompletos.
- Canciones escritas en idiomas distintos al español.

Normalización y Corrección Ortográfica. Para mejorar la coherencia del corpus y reducir la variabilidad léxica, se implementaron los siguientes pasos:

- Conversión a minúsculas para uniformar los datos y evitar diferencias artificiales entre palabras.
- Eliminación de signos de puntuación y caracteres especiales, a excepción de aquellos relevantes para la estructura gramatical.
- Corrección ortográfica automática utilizando la biblioteca pypellchecker [15], con un diccionario de palabras en español, asegurando que las palabras estuvieran correctamente escritas.

Estandarización de Expresiones. Dado que las letras de los corridos suelen incluir palabras coloquiales o variantes ortográficas propias del género, se compararon las palabras del corpus con un diccionario de términos en español para detectar y corregir términos no estándar o con errores tipográficos frecuentes.

El corpus final se dividió en dos conjuntos de datos:

Tabla 1. Características de los Datasets.

Género	Número total de canciones	Número de artistas únicos
Corridos	2,496	300
Regional	2,944	286

Pseudocódigo 1. Extracción de entidades.

```

FUNCIÓN obtener_terminos_ner(texto)
  doc ← aplicar_modelo_NER(texto) // Procesa el texto con el modelo NER
  terminos ← lista_vacia // Lista para almacenar los términos extraídos
  PARA cada entidad ent EN doc.entidades
    SI contar_palabras(ent.texto) > 1 ENTONCES
      agregar_a_lista(terminos, convertir_a_minusculas(ent.texto))
    FIN SI
  FIN PARA
  RETORNAR terminos
FIN FUNCIÓN

```

Estas transformaciones permitieron construir un conjunto de datos limpio y estructurado, optimizando la detección de patrones temáticos en los corridos y la música regional.

c. Extracción de términos con NER

Para obtener los términos clave, se implementó la función “obtener_terminos_ner”, con el objetivo de procesar el texto utilizando el modelo “es_core_news_md” de la biblioteca de “spaCy” y extraer las entidades compuestas por dos o más palabras. Se presenta su estructura en el **Pseudocódigo 1** para facilitar la comprensión del proceso.

Esta función primero procesa el texto con el modelo de spaCy para realizar el Reconocimiento de Entidades Nombradas (NER), luego filtra las entidades detectadas asegurando que solo se incluyan aquellas con más de una palabra, y finalmente las convierte a minúsculas antes de devolver la lista de términos extraídos. Este enfoque permite identificar palabras clave que reflejan la estructura y el contenido del dataset analizado.

d. Análisis de frecuencia de términos

Tras la extracción de términos **NER**, se aplicó un análisis de frecuencia para identificar las palabras más comunes dentro del dataset. Para ello, se utilizó el **Pseudocódigo 2**.

Pseudocódigo 2. Análisis de frecuencia.

```
INICIO
// Descomponer listas en filas individuales
terminos ← descomponer_listas(df['terminos'])
// Contar la frecuencia de cada término
conteo ← contar_frecuencia(terminos)
// Convertir conteo a una tabla para mejor visualización
df_frecuencia ← crear_tabla(conteo, columnas=['Término', 'Frecuencia'])
// Ordenar la tabla por frecuencia en orden descendente
df_frecuencia ← ordenar_por_columna(df_frecuencia, 'Frecuencia', descendente=True)
// Mostrar los 50 términos más frecuentes
imprimir(primeros_n_elementos(df_frecuencia, 50))
FIN
```

Este procedimiento permite descomponer las listas de términos en filas individuales, contar su frecuencia y estructurar los datos en una tabla ordenada. Posteriormente, se filtran los términos con una frecuencia mínima de 1 y se seleccionan los 50 más comunes para su análisis.

e. Preprocesamiento de datos para detección de tópicos

BERTopic, desarrollado por Grootendorst [13], es un modelo de detección de tópicos basado en técnicas de aprendizaje profundo que combina embeddings generados por modelos de lenguaje preentrenados con métodos de reducción de dimensionalidad y clustering.

Su procedimiento consta de cuatro etapas principales:

Generación de incrustaciones de texto. Se utilizan modelos de lenguaje preentrenados para transformar cada documento (letra de canción) en una representación vectorial en el espacio semántico. Para este estudio, se empleó el modelo MiniLM-L6-V2, optimizado para tareas de representación de textos.

Reducción de dimensionalidad. Dado que los embeddings suelen tener una dimensionalidad alta, se aplica UMAP (Uniform Manifold Approximation and Projection) para reducir la complejidad del espacio y preservar relaciones semánticas clave entre documentos.

Agrupamiento de documentos. Los embeddings reducidos son agrupados mediante HDBSCAN (Hierarchical Density-Based Spatial Clustering of Applications with Noise), lo que permite identificar grupos de canciones con características temáticas similares. Este enfoque permite gestionar la variabilidad de los textos sin necesidad de especificar un número fijo de tópicos.

Extracción de términos representativos. Para obtener una representación semántica interpretable de cada tema, se utiliza c-TF-IDF (class-based Term Frequency-Inverse Document Frequency), que permite identificar las palabras más relevantes dentro de cada clúster, diferenciándolas de otros tópicos en el datasets.

f. Análisis de sentimientos

El modelo utilizado para realizar el análisis de sentimientos fue `edumunozsala/roberta_bne_sentiment_analysis_es`. Este modelo se basa en la arquitectura RoBERTa-base-bne, una variante de RoBERTa preentrenada con el mayor corpus en español conocido hasta la fecha, con un total de 570 GB. Este corpus fue desarrollado por la Biblioteca Nacional de España [14].

El modelo fue afinado específicamente para análisis de sentimientos en español utilizando un conjunto de datos de aproximadamente 50,000 reseñas de películas en español. Este conjunto de datos está equilibrado y proporciona cada reseña en inglés y español, junto con las etiquetas en ambos idiomas.

g. Análisis de emociones

Para el análisis de emociones se utilizó el modelo `bhadresh-savani/distilbert-base-uncased-emotion`, el cual es una versión de DistilBERT, un modelo de lenguaje que utiliza destilación de conocimiento durante la fase de preentrenamiento para reducir el tamaño de BERT en un 40%, manteniendo el 97% de su capacidad de comprensión del lenguaje. Este modelo es más pequeño y rápido que BERT y otros modelos basados en BERT.

La arquitectura de DistilBERT se basa en la de BERT, pero con algunas modificaciones para hacerlo más eficiente.

El modelo `distilbert-base-uncased-emotion` ha sido ajustado específicamente para la clasificación de emociones utilizando el conjunto de datos de emociones de Twitter.

h. Análisis de intención

Para este análisis se empleó BART, el cual es un modelo de transformador secuencia a secuencia que combina las características de un codificador bidireccional, similar a BERT, y un decodificador autorregresivo, similar a GPT. Esta arquitectura le permite manejar tareas de generación y comprensión de lenguaje natural de manera eficiente. El proceso de preentrenamiento de BART implica dos etapas principales:

1. Corrupción del texto: Se aplica una función de ruido al texto original para corromperlo, lo que puede incluir la eliminación de tokens, permutaciones de oraciones, entre otros métodos.

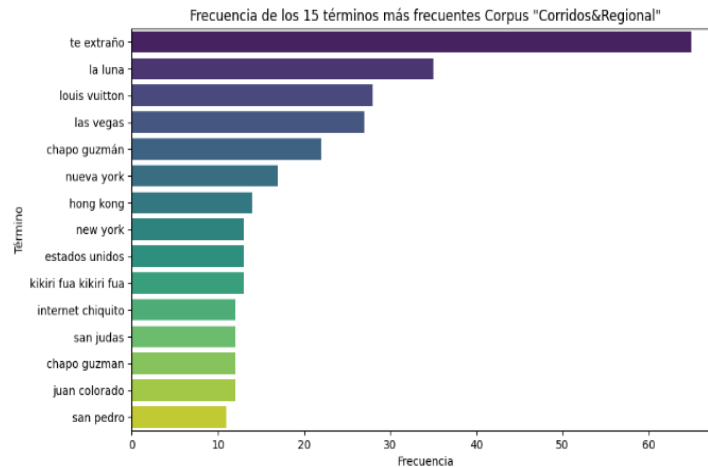


Fig. 1. Los 15 términos más frecuentes en el dataset "Corridos&Regional".

2. Reconstrucción del texto: El modelo aprende a reconstruir el texto original a partir de la versión corrompida, desarrollando así una comprensión profunda de la estructura y el significado del lenguaje.

Para este trabajo se especificaron las etiquetas a detectar por el modelo: “Glorificar el narcotráfico”, “Describir el narcotráfico” y “Desprestigiar el narcotráfico”.

4. Resultados y experimentos

En esta sección se muestran los resultados de los experimentos realizados en cada modelo previamente mencionado.

a. Extracción de términos con NER

Para este experimento se realizó el análisis de los dos datasets de música regional y corridos y un tercer dataset que es la fusión de estos dos últimos. Aquí los resultados:

Dataset fusión “corridos y regional”. En este dataset se analizaron 5141 letras de canciones, los resultados de los términos más comunes de todo el listado de las canciones se muestran en la **Fig. 1**.

Distribución de términos y repetición de entidades:

- El término “te extraño” destaca significativamente, esto sugiere una presencia importante de temas románticos. Varios términos hacen referencia a lugares como: “las vegas”, “nueva york”, “Hong kong” y “estados unidos”. En corridos y música regional es común aludir a ciudades de EE.UU. donde hay comunidades mexicanas o bien actividades delictivas y de negocios.
- Aparecen nombres propios ligados a la cultura del narcotráfico, por ejemplo: “chapo guzmán”, “juan colorado”.

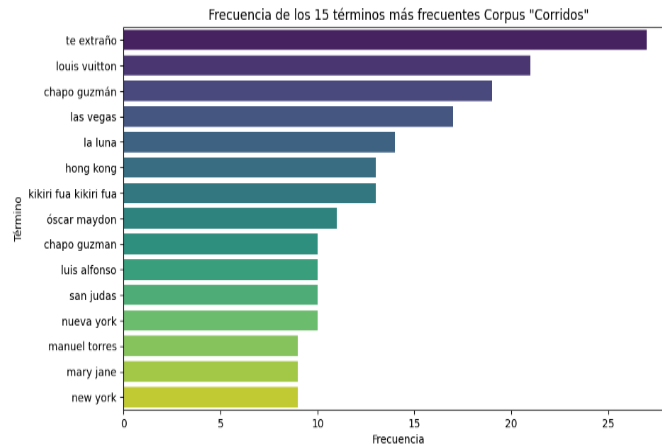


Fig. 2. Los 15 términos más frecuentes en el dataset "Corridos".

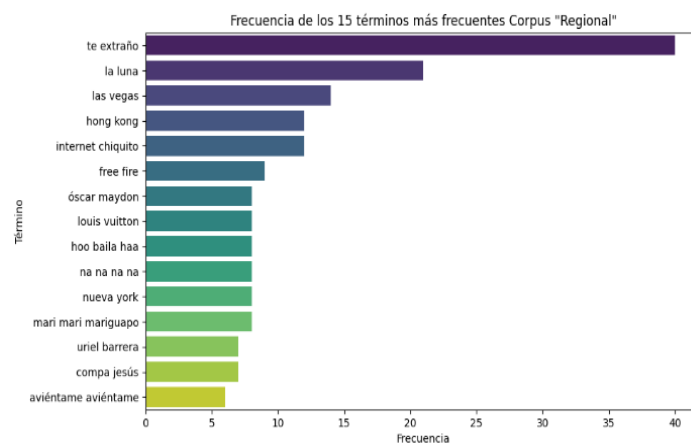


Fig. 3. Los 15 términos más frecuentes en el dataset "Regional".

Dataset “Corridos”. En este dataset se analizaron 2,496 letras de canciones, los resultados de los términos más comunes de todo el listado de las canciones se muestran en la **Fig. 2**.

Distribución de términos y repetición de entidades:

- Al igual que en la Fig. 1, el término “te extraño” vuelve a aparecer como el termino con mayor frecuencia. Esto refuerza la idea de que muchos corridos modernos también hay una presencia de temas románticos, más allá de los tradicionales temas de narcotráfico, migración o aventuras.
- El término “louis vuitton” ocupa el segundo lugar. En corridos actuales, es frecuente la mención de marcas de lujo para enfatizar con el poder adquisitivo o el estatus del personaje.

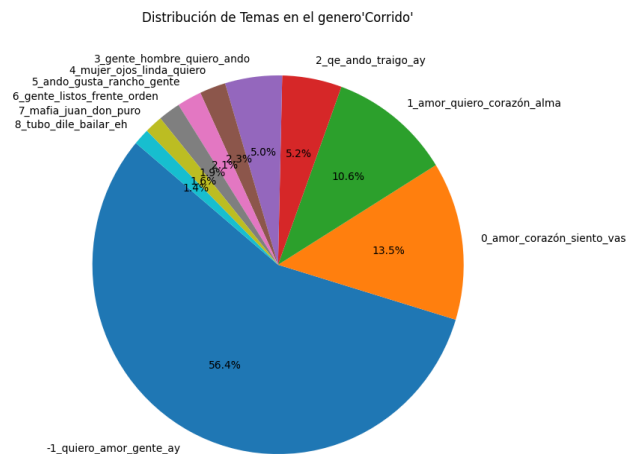


Fig. 4. Distribución de los 10 temas principales del dataset del género 'Corrido'.

Dataset “Regional”. En este dataset se analizaron 2,944 letras de canciones, los resultados de los términos más comunes de todo el listado de las canciones se muestran en la **Fig. 3**.

Distribución de términos y repetición de entidades:

- Nuevamente “te extraño” es el término más frecuente, lo que sugiere que la música regional mexicana lo temas de amor y desamor son muy comunes.
- Se menciona “las vegas”, “hong kong” y “nueva york”, lo que sugiere la influencia de la migración o la conexión con la vida en el extranjero.

b. Modelado de tópicos con BERTopic

Se realizó el análisis temático utilizando el modelo BERTopic, basado en técnicas avanzadas de embeddings y clustering semántico. Los resultados del análisis mostraron diferencias notables en la distribución de tópicos entre los géneros musicales analizados. En el género "Corrido", destacaron tópicos relacionados con emociones personales y afectivas, como "amor_corazón_siento_vas" (13.5%) y "amor_quiero_corazón_alma" (10.6%). Un tema particularmente relevante fue "mafia_juan_don_puro", este tema sugiere que diversas canciones del dataset comparten un contexto común o una combinación temática en la que dichos términos son frecuentes y distintivos. Sin embargo, esto no implica necesariamente que se refieran a un único personaje llamado Juan. Más bien, podría indicar que términos como 'don' (empleado como título respetuoso o jerárquico), 'mafia' (asociado al crimen organizado o contextos ilícitos) y nombres propios como 'Juan', aparecen regularmente juntos en letras que posiblemente reflejan narrativas comunes en los corridos, relacionadas con figuras de poder o historias del narcotráfico. (véase la **Fig. 4**).

Por otro lado, en el género "Regional", los tópicos más frecuentes fueron "amor_quiero_vida_corazón" (33.9%) y "tatara_compa_tatara" (6.2%), resaltando una predominancia en temas relacionados con aspectos emocionales, la vida cotidiana y las

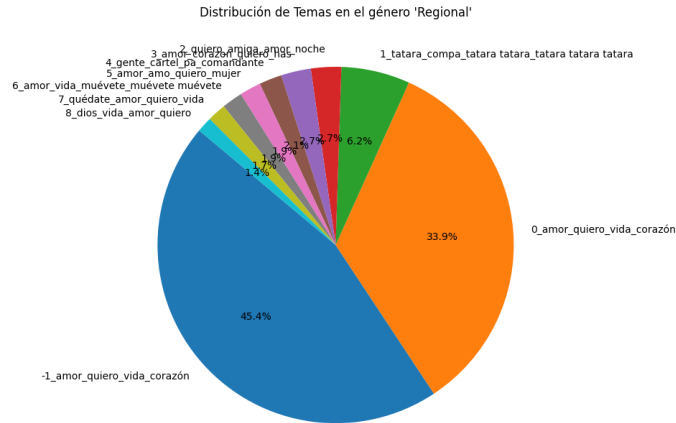


Fig. 5. Distribución de los 10 temas principales del dataset del género 'Regional'.

Tabla 2. Cantidad de sentimientos en las letras de los corridos y la música regional.

Etiqueta	Corridos	%	Regional	%
Negativo	440	17.63%	1227	41.68%
Neutral	2056	82.37%	1717	58.32%
Positivo	0	0%	0	0%

relaciones personales (véase **Fig. 5.**). Este análisis proporciona una visión clara sobre las temáticas predominantes y evidencia tanto diferencias como similitudes en las narrativas entre ambos géneros.

c. Análisis de sentimientos

El modelo `edumunozsala/roberta_bne_sentiment_analysis_es` fue entrenado utilizando Amazon SageMaker y el contenedor de Deep Learning de Hugging Face. En términos de rendimiento, el modelo alcanzó una precisión del 91.07%, una puntuación F1 de 90.91%, una precisión (*precision*) de 90.64% y una cobertura (*recall*) de 91.18% en el conjunto de datos de prueba. Este modelo es adecuado para tareas de análisis de sentimientos en textos en español, fue aplicado directamente a nuestro dataset ya que el modelo proporciona indicadores útiles sobre la distribución de sentimientos en las letras, estos deben interpretarse como aproximaciones del modelo.

Aunque no existen estudios previos que apliquen exactamente esta combinación de técnicas a los corridos, nuestra metodología se inspira en trabajos que han aplicado modelos similares a otros dominios. Por ejemplo, [9] utilizó RoBERTa para análisis de sentimientos en letras de canciones, mientras que [10] demostró la efectividad de BART-large-MNLI para clasificación de intención en diversos tipos de texto. Nuestra contribución radica en adaptar estos enfoques al dominio específico de la música regional mexicana.

Tabla 3. Cantidad de emociones en las letras de los corridos y la música regional.

Etiqueta	Corridos	%	Regional	%
Tristeza	1048	41.99%	1236	41.98%
Miedo	474	18.99%	559	18.99%
Enojo	300	12.02%	353	11.99%
Alegría	349	13.98%	412	13.99%
Sorpresa	200	8.01%	236	8.02%
Amor	125	5.01%	147	4.99%

Véase en la **Tabla 2** como se representa la frecuencia de ciertos sentimientos en las letras de estas canciones que proporcionan una visión general de cómo se expresan los sentimientos en ambos géneros de música.

d. Análisis de emociones

El modelo **bhadresh-savani/distilbert-base-uncased-emotion**, el cual es una versión de DistilBERT, alcanzó una precisión de 93.8% y una puntuación F1 de 93.79% en el conjunto de datos de emociones de Twitter.

En la **Tabla 3** se representa la frecuencia de ciertas emociones en las letras de estas canciones. Se enfoca en la combinación de ambos géneros (corridos y regional), por lo que proporciona una visión general de cómo se expresan las emociones en ambos géneros de música.

El análisis de 2,944 canciones de música regional y 2,496 corridos revela una fuerte presencia de emociones negativas, especialmente tristeza (42%), miedo (19%) y enojo (12%), reflejando temáticas de dolor, peligro y protesta. Las emociones positivas, como alegría (14%) y amor (5%), están menos representadas, y la sorpresa (8%) aparece ligada a giros narrativos.

En cuanto a sentimientos, predomina un tono neutral: en la música regional con 58%, y en los corridos con 82%, lo que sugiere un estilo más descriptivo o narrativo. Los sentimientos negativos son más notables en la música regional (42%) que en los corridos (18%), y los positivos son casi nulos (0%) en ambos géneros.

e. Detección de intención

Utilizando el modelo BART-large-MNLI y colocando las 3 etiquetas: Glorificar, Describir y Desprestigiar al narcotráfico se realizó el análisis de todas las canciones del repositorio dividido entre las canciones que se catalogaron en corridos y música regional. De esta forma el modelo puede indicar cual es la intención de la letra al menos con el enfoque semántico y contextual que ofrece la inferencia textual del modelo. Después del procesamiento del modelo con todas las canciones se obtuvieron los siguientes resultados:

Como se muestra en la **Tabla 4**, los porcentajes presentados corresponden al promedio de confianza en la consistencia de los resultados, calculado por el propio

Tabla 4. Cantidad de veces que se detectó cada etiqueta de la intención de la canción dentro de cada dataset.

Etiqueta	Corridos	%Conf.	Regional	%Conf.
Glorificar el narcotráfico	81	39%	93	39%
Describir el narcotráfico	1900	45%	2180	44%
Desprestigiar el narcotráfico	515	42%	671	42%

modelo. Aunque estos valores son relativamente bajos, esto se debe a que el modelo consideró las canciones en su totalidad para realizar la inferencia.

También hay que tomar en cuenta otros factores que pueden afectar este análisis general de las canciones. Porque una canción puede solo narrar (describir) cosas referentes al narcotráfico, pero hacerlo desde una perspectiva de protagonista lo que puede llevar a que se describan historias, aspectos o situaciones sobre el narcotráfico de forma positiva no directamente, afectando el análisis por parte del modelo.

5. Conclusiones y trabajo futuro

El análisis de los corridos y la música regional mexicana ha permitido identificar patrones lingüísticos, temáticos, emocionales e intención que reflejan la evolución del género y su impacto en la sociedad. A través del **Reconocimiento de Entidades Nombradas (NER)** y el modelado de tópicos con **BERTopic**, se han detectado referencias recurrentes a emociones personales, lugares internacionales y figuras del crimen organizado, evidenciando la coexistencia de temas afectivos con narrativas sobre narcotráfico y globalización.

El análisis de sentimientos muestra un tono predominantemente negativo, con una ausencia de elementos festivos o positivos, lo que refuerza la idea de los corridos como crónicas de violencia, sacrificio y resistencia cultural. Además, la clasificación de intención con **BART-large-MNLI** resalta la necesidad de diferenciar entre canciones que glorifican el crimen y aquellas que lo narran de manera neutral.

Como líneas futuras de investigación, se propone extender el estudio al análisis transversal del lenguaje en este género musical, explorando cómo han evolucionado las expresiones, temas y referentes culturales a lo largo del tiempo, evidenciando la adaptación de estos géneros a las transformaciones sociales. Además, con el objetivo de mejorar la precisión en las inferencias, se plantea la implementación de un etiquetado a nivel de renglón. Esta aproximación permitiría promediar los valores obtenidos en cada segmento, brindando así un análisis más detallado y fino sobre el contenido y la carga emocional de las canciones.

Agradecimientos. Se agradece el apoyo del Tecnológico Nacional de México campus CENIDET y a la SECIHTI a través de los registros de becas (CVU 1343321, CVU 1105435, CVU 1315217)

Divulgación de intereses. El autor José Eder Martínez Martínez ha recibido apoyo financiero mediante las becas CVU 1343321. El autor Yessenia Díaz Álvarez ha recibido apoyo financiero mediante las becas CVU 1105435. El autor Tlaloc Humberto Vergara Morales ha recibido apoyo financiero mediante las becas CVU 1315217. Los autores Víctor Manuel Millán Aguilar, Gabriel

González Serna y Noé Alejandro Castro Sánchez declaran no tener conflictos de intereses ni apoyos externos relacionados con esta publicación.

Referencias

1. Zhang, Y., Wang, X., Liu, Z.: Advances in Machine Learning for Natural Language Processing. In: J. Artif. Intell. Res., 45(2), pp. 123–145 (2023) doi: doi.org/10.1145/3604931.
2. Brants, T., Franz, A.: Web 1T 5-gram Corpus Version 1.1. In: Proc. Conf. Empirical Methods in Natural Language Processing (EMNLP 2011), Association for Computational Linguistics, pp. 123–130 (2011) url: <https://aclanthology.org/D11-1141.pdf>
3. Bojar, O., Chatterjee, R., Federmann, C.: Findings of the 2016 Conference on Machine Translation (WMT16). In: Proc. First Conf. Machine Translation, Association for Computational Linguistics, pp. 131–198 (2016) url: <https://aclanthology.org/S16-1002.pdf>
4. Fell, M., Cabrio, E., Tikat, M.: The Wasabi Song Corpus and Knowledge Graph for Music Lyrics Analysis. Lang. Resour. Eval., 57(1), pp. 89–119 (2023).
5. Kleedorfer, F., Knees, P., Pohle, T.: Oh oh oh whoah! Towards Automatic Topic Detection in Song Lyrics. In: Proc. ISMIR 2008, pp. 287–292 (2008).
6. Wickham, E.N.: Girlbosses, the Red Pill, and the Anomie and Fatale of Gender Online: Analyzing Posts from r/suicidewatch on Reddit. In: Proc. 3rd Int. Conf. Natural Language Processing for Digital Humanities and 8th Int. Workshop on Computational Linguistics for Uralic Languages, pp. 195–212 (2023)
7. Díaz-Santana, M.: Impacto sociocultural de los narcocorridos. Rev. Estud. Mexicanos, 22(3), pp. 34–56 (2014).
8. Serrano, A.: Los narcocorridos: Entre la apología del delito y la crítica social. Rev. Mex. Sociol., 71(1), pp. 51–74 (2009).
9. González-Ruiz, J.: Análisis de sentimiento en letras de canciones usando procesamiento de lenguaje natural. Acta Hispánica, 10(2), pp. 123–140 (2020).
10. Biswas, R.: Flask API for Hugging Face Facebook BART-large-MNLI. GitHub (2021) url: <https://github.com/rajatdiptabiswas/flask-api-hugging-face-fb-bart-large-mnli>
11. Dataloop AI: BART-Large-MNLI-Yahoo-Answers Model Overview. Dataloop AI (2023) url: https://dataloop.ai/library/model/joeddav_bart-large-mnli-yahoo-answers
12. Stack Overflow: How Does Hugging Face’s Zero-Shot Classification Work in A Production Web App? Stack Overflow (2023) url: <https://stackoverflow.com/questions/75866093>
13. Grootendorst, M.: Bertopic: Neural Topic Modeling with A Class-Based TF-IDF Procedure. (2022) arXiv:2203.05794.
14. Biblioteca Nacional de España (BNE): Corpus de Referencia del Español Actual (CREA). url:<https://www.bne.es/es/Catalogos/Corpus>
15. Chollet, F.: Keras. (2015) url: <https://keras.io>

Caracterización de la fatiga mental y su recuperación en tarea de conducción mediante señales electroencefalográficas utilizando algoritmos basados en fractales

Diana G. González-Rodríguez¹, Dulce Martínez-Peon^{1,2},
Erika L. Guzmán Arriaga³, Michelle Contró Esparza³,
Virginia García Pinedo⁴, Carlos A. Rivera Vergara⁴,
Gerardo Benavides Bravo⁵, Manuel A. Ramírez Hernández⁵.

¹ Departamento de Ingeniería Eléctrica y Electrónica,
Instituto Tecnológico de Nuevo León,
México

² División de Estudios de Posgrado e Investigación,
Instituto Tecnológico de Nuevo León,
México

³ Departamento de Ingeniería en Gestión Empresarial,
Instituto Tecnológico de Nuevo León,
México

⁴ Departamento de Ingeniería Industrial,
Instituto Tecnológico de Nuevo León,
México

⁵ Departamento de Ciencias Básicas,
Instituto Tecnológico de Nuevo León,
México

dulce.mp@nuevoleon.tecnm.mx

Resumen. Este estudio tiene como objetivo identificar la fatiga mental y su recuperación en conductores mediante señales electroencefalográficas (EEG). Se utilizó una base de datos de EEG públicamente disponible que contenía 60 minutos de datos de tareas de conducción de siete participantes. Las señales se dividieron en cuatro bloques a los 2, 25, 40 y 60 minutos para su análisis. El preprocesamiento incluyó técnicas de filtrado y análisis de componentes independientes (ICA), seguido de análisis wavelet. Se emplearon tres algoritmos para la caracterización de la señal: desincronización/sincronización relacionada con eventos (ERD/ERS), el exponente de Hurst y la dimensión fractal de Higuchi, centrándose en las bandas alfa y theta. Los resultados de ERD/ERS indican que la fatiga mental puede detectarse después de 25 minutos de conducción continua, con signos de recuperación observados 15 minutos después. Además, se encontró que las variaciones en las señales de EEG dentro de las regiones frontocentral y parietal se correlacionan con la aparición de la fatiga. De manera similar, los algoritmos de Hurst y Higuchi produjeron resultados comparables, lo que refuerza el potencial de la monitorización basada en EEG para la detección temprana de la fatiga mental. Estos hallazgos brindan

información valiosa para el desarrollo de intervenciones de seguridad en entornos de conducción.

Palabras clave: Fatiga mental, señales EEG, exponentes de Hurst, dimensión fractal de Higuchi.

Characterization of Mental Fatigue and its Recovery during a Driving Task Using Electroencephalographic Signals and Fractal-based Algorithms

Abstract. This study aims to identify mental fatigue and its recovery in drivers using electroencephalographic (EEG) signals. A publicly available EEG database containing 60 minutes of driving task data from seven participants was used. The signals were divided into four blocks at 2, 25, 40, and 60 minutes for analysis. Preprocessing included filtering techniques and independent component analysis (ICA), followed by wavelet analysis. Three algorithms were employed for signal characterization: event-related desynchronization/synchronization (ERD/ERS), the Hurst exponent, and Higuchi's fractal dimension, focusing on the alpha and theta bands. The ERD/ERS results indicate that mental fatigue can be detected after 25 minutes of continuous driving, with signs of recovery observed 15 minutes later. Additionally, variations in EEG signals within the frontocentral and parietal regions were found to correlate with the onset of fatigue. Similarly, the Hurst and Higuchi algorithms produced comparable results, reinforcing the potential of EEG-based monitoring for the early detection of mental fatigue. These findings provide valuable insights for the development of safety interventions in driving environments.

Keywords: Mental fatigue, EEG signals, Hurst exponents, Higuchi fractal dimension.

1. Introducción

Actualmente, el entorno laboral se caracteriza por una alta tasa de horas trabajadas, lo que ha llevado a un aumento de la fatiga mental [1,2]. Este fenómeno, reconocido como un importante desafío de salud pública, no solo afecta el bienestar emocional de los trabajadores, sino que además tiene consecuencias significativas para la productividad y el rendimiento del trabajador [2,3]. En este contexto, diversos estudios científicos respaldan la necesidad de comprender y abordar la fatiga mental de manera efectiva. La relación entre la fatiga mental y la productividad ha sido objeto de varias investigaciones. Por otro lado, se ha demostrado que la fatiga mental es uno de los principales causantes de accidentes vehiculares a nivel mundial; cada año, las colisiones causadas por el tránsito cobran la vida de aproximadamente 1,19 millones de personas, según informes de la OMS. Es importante mencionar que estudios de la misma organización indican que la fatiga mental representa cerca del 30% de los accidentes de tránsito [4]. La fatiga mental se puede medir a través de baterías psicológicas y en los años recientes se han desarrollado dispositivos que usan

indicadores fisiológicos como la presión sanguínea, frecuencia cardíaca o parpadeo, o basados en señales eléctricas como el electrooculograma o la electroencefalografía (EEG) [5,6]. Sin embargo, hasta ahora no se ha analizado cómo se presenta la recuperación luego de la fatiga.

Las señales EEG requieren de un procesamiento para poder encontrar las características en tiempo y frecuencia de la tarea o enfermedad deseada. Cuando se presenta la fatiga mental, en la región frontocentral se presenta a los 25 minutos una desincronización (decremento de la potencia espectral) en la banda alfa (8–12 Hz), y una sincronización (aumento de la potencia espectral) en la región parietal, [6-10]. Además, otros estudios analizaron la banda theta para detectar la fatiga mental, [11-12]. Esto se ha demostrado utilizando la técnica de desincronización/sincronización relacionada con eventos (ERD/ERS) [13].

Está técnica ha sido utilizada desde hace mucho tiempo pero tiene la desventaja de que para su procesamiento se requieren ventanas de tiempo de 3 segundos, por lo que un análisis en tiempo real no es factible. Existen técnicas para caracterizar que utilizan ventanas de tiempo de 100 ms a 300 ms, entre las que se encuentran el índice de Hurst o Higuchi.

Por otro lado, la naturaleza de las señales EEG es que son señales no regulares, no estacionarias, tienen características fractales y se ha demostrado que pueden ser caracterizadas mediante algoritmos basados en fractales como el exponente de Hurst, Wavelets, Dimensión Fractal de Higuchi (HFD) [14-15].

El propósito de este proyecto es caracterizar señales EEG mediante algoritmos basados en fractales como Higuchi y Hurst para detectar fatiga mental en sujetos sanos en tareas de conducción de vehículos, así como la respuesta ante la recuperación después de que se expresa la fatiga y sin dejar de realizar la tarea, para ello, se analizaron señales de una base de datos pública.

Para el análisis de las señales se generaron 4 bloques con ventanas de tiempo de 2, 25, 40 y 60 minutos, se aplicaron métricas como precisión y exactitud, con el fin de evaluar su eficacia en la detección de la fatiga mental de manera más precisa y robusta.

Por último, se generó un análisis ANOVA para la banda alfa y la banda theta, analizando los electrodos FC3, FC4, P3, P4, T3 y T4, donde se realizó la comparación entre los 4 bloques y demostrar que existe diferencia significativa en el bloque 1-2 y 3-4.

2. Materiales y métodos

En esta sección se presentan los materiales utilizados, así como la metodología implementada.

2.1 Materiales

Hardware y software. Se utilizó una computadora Lenovo™ IdeaPad330s con Windows 10 de 64 bits, procesador Intel Core i5 fe 8va generación, 12 GB de RAM, se procesaron los datos en MATLAB® R2022b.

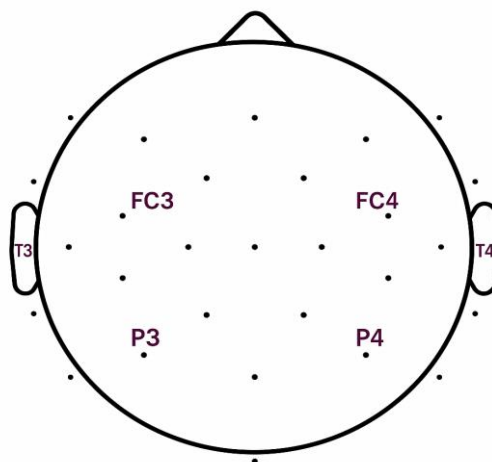


Fig. 1. Distribución de los electrodos de EEG en el sistema 10-20. Los electrodos seleccionados para el análisis son FC3, FC4, P3, P4, T3 y T4.

Base de datos. Se utilizó una base de datos en línea utilizando 7 de los 27 participantes registrados (25 ± 3 años), estos 7 sujetos fueron seleccionados porque son los que tienen un registro completo de 60 minutos, lo cual permite hacer un análisis en ese rango de tiempo. En el estudio se utilizaron 30 electrodos de EEG, de los cuales se analizaron solamente 6 ubicados en las zonas fronto-central (FC), temporal (T) y parietal (P), ver Figura 1.

El experimento consistió en llevar a cabo una simulación de manejo en el laboratorio de la universidad, el entorno donde se desarrolló la conducción fue en una carretera recta, de forma aleatoria se presentaron eventos (desviaciones) para medir el tiempo que le tomaba al participante en responder, denominaron los eventos como 251, 252, 253 y 254. En el evento 251 se mostraba una desviación hacia la izquierda, en el 252 una desviación a la derecha, en el evento 253 se mostraba el tiempo que tardó la persona en responder para posicionar el automóvil de manera correcta dentro de la carretera y por último en el evento 254 se mostraba el tiempo total desde el inicio hasta el final de cada evento. La finalidad del experimento fue evaluar el deterioro de la atención sostenida y el momento donde se presenta la fatiga mental [16].

2.2 Métodos

El análisis de señales EEG, se divide en 3 etapas, la primera consiste en el preprocesamiento de las señales, donde se agruparon los datos, se realizó la reducción del muestreo y se aplicaron filtros (pasa altas, pasa bajas y paso banda), la segunda etapa consistió en recortar las señales por épocas, se analizaron por componentes independientes (ICA) y análisis de Wavelet. En la tercera etapa las señales se caracterizaron utilizando la técnica de ERD/ERS, Higuchi y el exponente de Hurst, ver Figura 2.

Preprocesamiento:

- 1) Agrupamiento de datos: se agruparon en cuatro bloques a 2, 25, 40 y 60 minutos.
- 2) Reducir muestreo: en esta etapa se reduce el número de muestras conservando la información, esto se realiza para reducir el cómputo de los datos. Se debe considerar que esta reducción no debe incumplir el teorema de Nyquist.
- 3) Filtro pasa altas: las señales por debajo de la frecuencia de corte (banda de frecuencia de corte) son atenuadas por el filtro de paso alto, mientras que las que están por encima de la frecuencia de corte (banda de frecuencia de paso) pasan.
- 4) Filtro pasa bajas: este filtro permite que pasen las señales por debajo de la frecuencia de corte (banda de paso) y atenúa las señales por encima de la frecuencia de corte (banda de corte).
- 5) Filtro pasa bandas: este filtro realiza la función de pasar una frecuencia específica en un ancho de banda específico y atenúa la frecuencia fuera del ancho de banda.

Procesamiento de señales:

- 1) Épocas: se seleccionó cada época con duración de 6 segundos respecto a los eventos 251 y 252.
- 2) Análisis de componentes independientes (ICA): se utilizó para desmezclar las señales utilizando la función `runic` de MATLAB®.
- 3) Wavelets: se descompusieron las señales en 5 niveles para seleccionar la banda de frecuencia alfa (8–12 Hz).

– Caracterización:

- 1) ERD/ERS: es la disminución porcentual de la potencia espectral y se calcula como $ERD\% = (N - P) / P * 100$, donde N es el número total de ensayos y P es el promedio, [17-18].
- 2) Exponente de Hurst (HRS): mide la persistencia de los sistemas a lo largo del tiempo, [19-20]. El Rango reescalado es uno de los más comunes para su cálculo y es el usado en este trabajo. El valor de HRS está entre 0 y 1. Si $0 < H < 0,5$ indica anti-persistencia, si $H = 0,5$ indica falta de correlación serial (ruido blanco gaussiano) y si $0,5 < H < 1$ indica persistencia en los datos.
- 3) Dimensión Fractal de Higuchi: es otro algoritmo basado en fractales, es no lineal y su valor se encuentra entre 1 y 2.

– Análisis estadístico:

Como método estadístico para la validación de los datos se utilizó la prueba ANOVA para encontrar diferencia significativa entre los 4 bloques, considerando una $\alpha < 0.05$ como significativa.

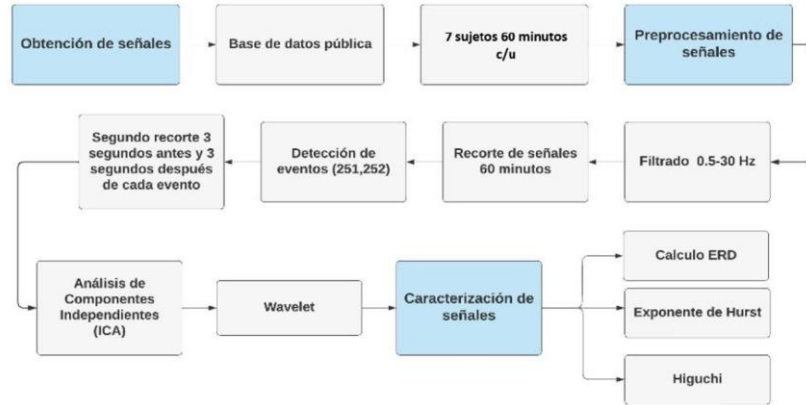


Fig. 2. Esquema del procedimiento para la detección de fatiga mental en conductores a partir de señales EEG. Se incluyen las etapas de obtención, preprocesamiento y caracterización de las señales, con técnicas como análisis de componentes independientes (ICA), wavelets y métricas fractales (ERD, Hurst y Higuchi).

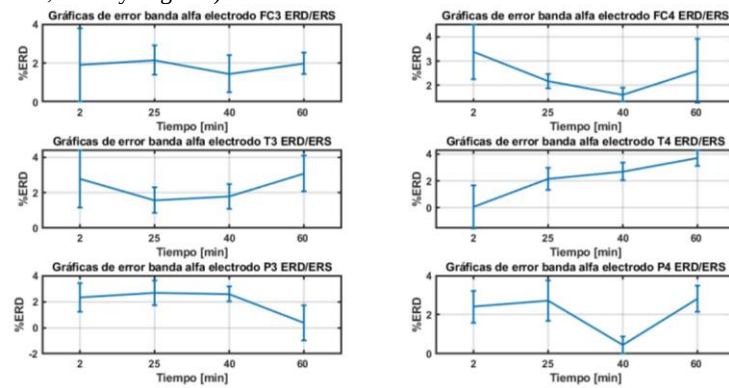


Fig. 3. Gráficas de error de banda alfa, electrodos FC3, FC4, T3, T4, P3 y P4 mediante la técnica ERD/ERS a lo largo de 60 minutos. El eje x muestra el tiempo de la tarea. El eje y muestra el porcentaje de respuesta ERD de la tarea.

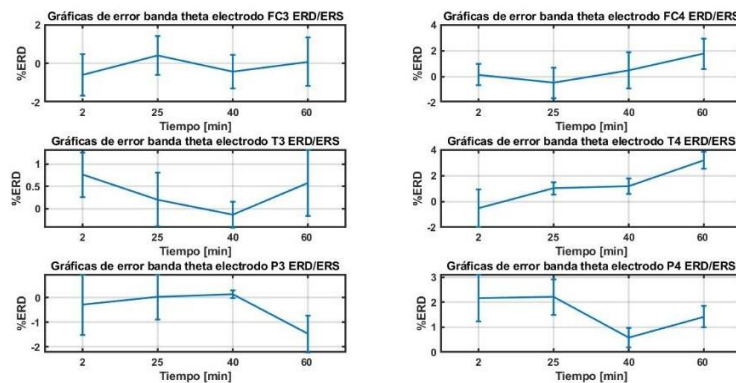


Fig. 4. Gráficas de error de banda theta, electrodos FC3, FC4, T3, T4, P3 y P4 mediante la técnica ERD/ERS a lo largo de 60 minutos.

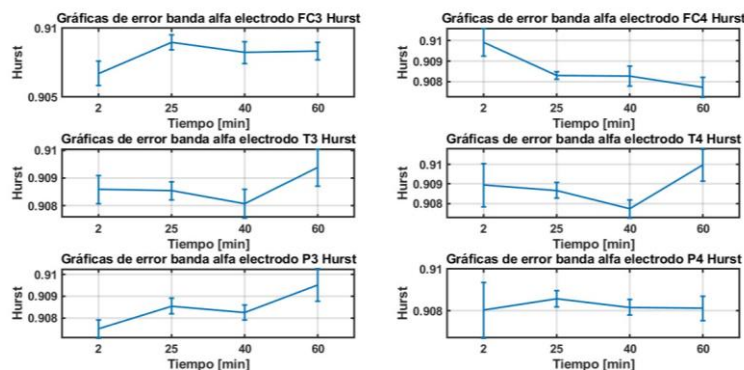


Fig. 5. Gráficas de error de banda alfa, electrodos FC3, FC4, T3, T4, P3 y P4 mediante la técnica de caracterización Hurst a lo largo de 60 minutos.

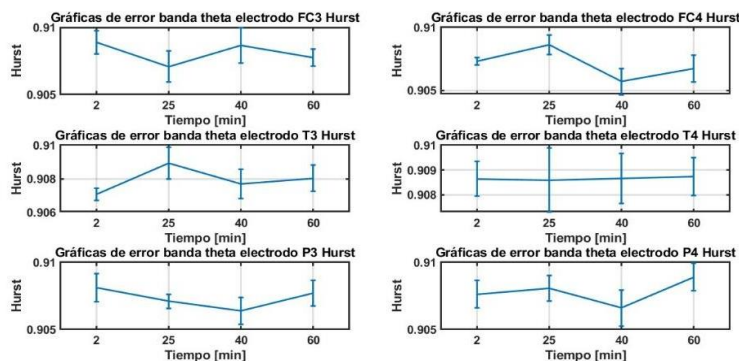


Fig. 6. Gráficas de error de banda theta, electrodos FC3, FC4, T3, T4, P3 y P4 mediante la técnica Hurst a lo largo de 60 minutos consecutivos.

3. Resultados

Los resultados obtenidos se presentan en las Figuras 3-8. Para el algoritmo con ERD/ERS se puede observar en la Figura 3 en la banda alfa que los electrodos T4 y P4 son los que muestran un cambio de amplitud entre los bloques 1 y 2 y posteriormente muestran otro cambio en el bloque 3.

En el caso de la banda theta, Figura 4, los electrodos FC3, FC4 y T4 muestran un cambio entre los bloques 1 y 2, y entre los bloques 2 y 3. En el caso del algoritmo con Hurst, los cambios entre los bloques 1, 2 y 3 en la banda alfa se muestran en los electrodos FC3, T3, P3 y P4. En el caso de la banda theta con Hurst se ven cambios entre los bloques 1, 2 y 3 en los electrodos FC3, T3, FC4 y P4. Finalmente, para Higuchi, los cambios entre bloques 1, 2 y 3 para la banda alfa se muestran en los electrodos FC4 y T4, mientras que para la banda theta estos cambios se presentan en los electrodos FC3, T3, P3, FC4 y P4.

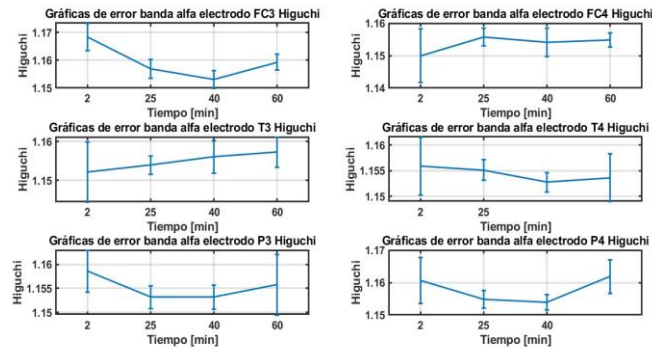


Fig. 7. Gráficas de error de banda alfa, electrodos FC3, FC4, T3, T4, P3 y P4 mediante la técnica Higuchi a lo largo de un período de 60 minutos.

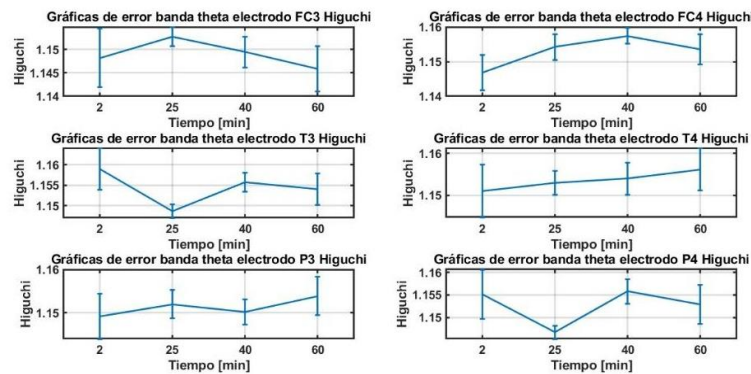


Fig. 8. Gráficas de error de banda theta, electrodos FC3, FC4, T3, T4, P3 y P4 mediante la técnica Higuchi en un período de 60 minutos consecutivos.

Tabla 1. Resultados del análisis de ANOVA para los tres algoritmos en la banda alfa.

ANOVA banda alfa						
Electrodo		Suma de cuadrados	DF	Media cuadrada	F	Valor p
FC3	(Intersección): Tiempo	15902	219	210.49	0.96772	0.61482
	Error (Tiempo)	2.86E+05	1314	217.51		
FC4	(Intersección): Tiempo	49084	219	224.13	1.0879	0.19794
	Error (Tiempo)	2.71E+05	1314	206.01		
P3	(Intersección): Tiempo	45675	219	208.56	1.0037	0.47602
	Error (Tiempo)	2.73E+05	1314	207.8		
P4	(Intersección): Tiempo	41740	219	190.59	1.0375	0.35092
	Error (Tiempo)	2.41E+05	1314	183.7		
T3	(Intersección): Tiempo	48631	219	22.06	1.0399	0.34274
	Error (Tiempo)	2.81E+05	1314	213.54		
T4	(Intersección): Tiempo	47009	219	214.65	1.0234	0.10178
	Error (Tiempo)	2.76E+05	1314	209.75		

Los resultados con ANOVA se muestran en la Tabla 1, en donde se puede observar que existe diferencia significativa entre los bloques 1, 2, 3 y 4 para los diferentes algoritmos empleados en la banda alfa.

Tabla 2. Resultados del análisis de ANOVA para los tres algoritmos en la banda theta.

ANOVAS banda theta						
Electrodo		Suma de cuadrados	DF	Media cuadrada	F	Valor p
FC3	(Intersección): Tiempo	40960	219	187.03	10.521	0.3019
	Error (Tiempo)	2.34E+05	1314	177.78		
FC4	(Intersección): Tiempo	43443	219	198.37	1.1335	0.10455
	Error (Tiempo)	2.30E+05	1314	175.01		
P3	(Intersección): Tiempo	34632	219	158.14	1.0422	0.33467
	Error (Tiempo)	1.99E+05	1314	151.73		
P4	(Intersección): Tiempo	35616	219	162.63	1.0035	0.47664
	Error (Tiempo)	2.13E+05	1314	162.06		
T3	(Intersección): Tiempo	33117	219	151.22	0.84963	0.93612
	Error (Tiempo)	2.34E+05	1314	177.98		
T4	(Intersección): Tiempo	43589	219	199.04	1.1801	0.048478
	Error (Tiempo)	2.22E+05	1314	168.66		

Los resultados con ANOVA se muestran en la Tabla 2, en donde se puede observar que existe diferencia significativa entre los bloques 1, 2, 3 y 4 para los diferentes algoritmos empleados en la banda theta.

Basado en el análisis realizado, se puede concluir que los datos presentados siguen una distribución normal. Esto se evidencia a través del gráfico de normalidad, el cual muestra que los puntos se alinean de manera cercana a la línea diagonal, lo que indica que la distribución de los datos se ajusta a una distribución normal ideal. Además, al examinar los resultados de la tabla ANOVA proporcionada, se observa que los p-valores de la prueba de normalidad de Shapiro-Wilk son en su mayoría superiores a 0.05, lo que no permite rechazar la hipótesis nula de normalidad.

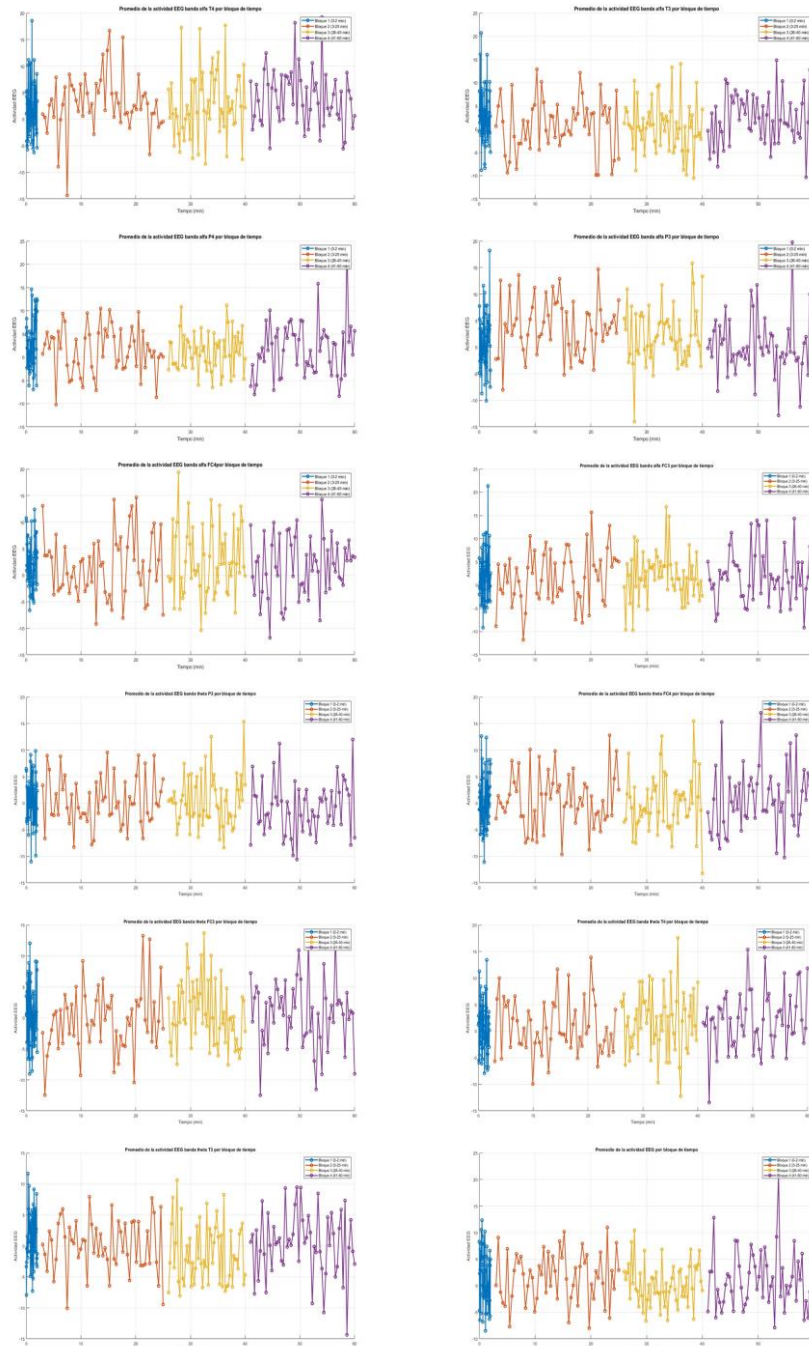
Complementariamente, al aplicar la prueba de Kolmogorov-Smirnov a los datos, se obtuvo un p-valor de 0.12, también mayor al nivel de significancia de 0.05, lo que respalda aún más la conclusión de que los conjuntos de datos analizados siguen una distribución normal. Los datos utilizados en este análisis fueron obtenidos a partir de los bloques de tiempo, posteriormente, se realizó la prueba de normalidad de Kolmogorov-Smirnov para cada sujeto, cuyos resultados confirmaron que los datos siguen una distribución normal, el análisis post hoc utilizando la prueba de Tukey-Kramer reveló diferencias significativas entre los grupos ($p < 0.05$) para la mayoría de los electrodos en ambas bandas de frecuencia.

4. Discusiones y conclusiones

En este estudio, se aplicaron algoritmos para analizar señales electroencefalográficas (EEG) con el objetivo de detectar la fatiga mental y su recuperación en una tarea de conducción vehicular, utilizando enfoques basados en fractales.

El análisis de los datos reveló un patrón cíclico en la aparición y recuperación de la fatiga mental, lo que indica que no es un estado permanente, sino un fenómeno dinámico que fluctúa con el tiempo.

Comprender esta dinámica es clave para desarrollar estrategias que optimicen el desempeño cognitivo y reduzcan los riesgos asociados con la conducción prolongada. Los resultados obtenidos ofrecen evidencia significativa sobre la evolución de la fatiga mental a lo largo del tiempo.



A través del análisis de ERD/ERS, el exponente de Hurst y la dimensión fractal de Higuchi, se identificó que los sujetos comenzaron a mostrar signos de fatiga tras 25

minutos de conducción continua (bloque 2), con signos de recuperación 15 minutos después (bloque 3). Este hallazgo es consistente con estudios previos que han demostrado que la fatiga mental surge rápidamente en tareas repetitivas y monótonas.

Cada técnica de análisis reveló cambios específicos en la actividad cerebral. La metodología ERD/ERS mostró que la disminución de las ondas alfa y el incremento de las ondas theta en los electrodos FC3, T3 y P4 reflejan un descenso en la capacidad de procesamiento sensoriomotor y atención sostenida. La reducción de las ondas alfa sugiere menor capacidad de control cognitivo, mientras que el aumento de las ondas theta indica una mayor predisposición a la somnolencia y disminución del estado de alerta, lo que podría afectar la respuesta del conductor ante situaciones imprevistas. Por otro lado, los análisis mediante Hurst muestra que puede ser utilizado para detectar la fatiga mental, como lo hace la técnica de ERD/ERS, esto debido a que la presenta un cambio en la respuesta del bloque 2 respecto al bloque 1. En la banda alfa, se observaron relaciones directas consistentes en los electrodos P3 y P4, donde tanto ERD/ERS como el exponente de Hurst presentaron una disminución hasta el minuto 40. Esta convergencia sugiere una modulación sincronizada de la actividad oscilatoria y de la persistencia fractal en regiones parietales ante la evolución de la fatiga mental. En un estudio anterior de los autores [21], se implementó el índice de Hurst a esta base de datos, sin embargo, solo se realizó un análisis a los 40 minutos, faltó un análisis acerca de la recuperación, la cual es mostrada en este estudio. Para el caso del análisis con Higuchi, no se muestran respuestas como Hurst, por lo que para trabajo a futuro se propone utilizar otras regiones cerebrales para analizar su respuesta.

La identificación temprana de estos cambios mediante el análisis de EEG podría ser clave para la implementación de sistemas de monitoreo de fatiga en tiempo real. Esto permitiría desarrollar intervenciones efectivas para mejorar la seguridad en la conducción y reducir el riesgo de accidentes relacionados con la fatiga mental. En conclusión, este estudio aporta conocimientos fundamentales sobre la dinámica de la fatiga mental y sus manifestaciones neurofisiológicas, ofreciendo una base sólida para futuras investigaciones enfocadas en su aplicación en entornos reales de conducción.

Referencias

1. Allison, P., Tiesman, H.M., Wong, I.S.: Working Hours, Sleep, and Fatigue in the Public Safety Sector: A Scoping Review of the Research. *Am J Ind Med.*, 65, pp. 878–897 (2022) doi: 10.1002/ajim.23407
2. Cunningham, T.R., Guerin, R.J., Ferguson, J.: Work-Related Fatigue: A Hazard for Workers Experiencing Disproportionate Occupational Risks. *Am J Ind Med.*, 65, pp. 913–925 (2022) doi:10.1002/ajim.23325
3. Gaitán-Domínguez, M., González-Rodríguez, M. A., Ordoñez-Argote, A. M.: ¿Cómo afecta la salud mental en la productividad laboral? Corporación Unificada Nacional de Educación Superior-CUN (2023) url: <https://repositorio.cun.edu.co/handle/cun/6084>
4. Organización Mundial de la Salud (OMS): Global Status Report on Road Safety (2023) url: <https://www.who.int/publications/i/item/9789241565684>
5. Dickens, A., Champion, A. J., Schenke, K. C.: Fatigue Management: a Systematic Review of Objective Measurement Techniques for Cognitive Fatigue. *Journal of Clinical and Experimental Neuropsychology*, 46(8), pp. 776–793 (2024).

6. Magnuson, J. R., Doesburg, S. M., McNeil, C. J.: Development and Recovery Time of Mental Fatigue and Its Impact on Motor Function. *Biological psychology*, 161 108076 (2021)
7. Lin, C. T., Wu, R. C., Liang, S. F.: EEG-Based Drowsiness Estimation for Safety Driving Using Independent Component Analysis. In: *IEEE Transactions on Circuits and Systems I: Regular Papers*, 52(12), pp. 2726–2738 (2005) doi: 10.1109/TCSI.2005.857555
8. Lin, C. T., Wu, R. C., Jung, T. P.: Estimating Driving Performance Based on EEG Spectrum Analysis. In: *EURASIP Journal on Advances in Signal Processing*, 2005 (19), pp. 1–10 (2005) doi:10.1155/ASP.2005.3165
9. Han, C., Sun, X., Yang, Y.: Brain Complex Network Characteristic Analysis of Fatigue During Simulated Driving Based on Electroencephalogram Signals. *Entropy*, vol. 21, no. 4, pp. 353 (2019) doi: 10.3390/e21040353
10. Gharagozlou, F., Saraji, G. N., Mazloumi, A.: Detecting Driver Mental Fatigue Based on EEG Alpha Power Changes During Simulated Driving. *Iranian Journal of Public health*, 44(12), pp. 1693 (2005) PMID: 26811821 PMID: PMC4724743.
11. Wascher, E., Rasch, B., Sanger, J.: Frontal Theta Activity Reflects Distinct Aspects of Mental Fatigue. *Biological psychology*, 96, pp. 57–65 (2014) doi: 10.1016/j.biopsycho.2013.11.010.
12. Goodman, S.P.J., Collins, B., Shorter, K.: Approaches to Inducing Mental Fatigue: A Systematic Review and Meta-Analysis Of (Neuro)Physiologic Indices. *Behav Res*, 57(102) (2025) doi: 10.3758/s13428-025-02620-7.
13. Dujardin, K., Derambure, P., Defebvre, L.: Evaluation of Event-Related Desynchronization (ERD) During A Recognition Task: Effect of Attention. *Electroencephalography and clinical neurophysiology*, 86(5), pp. 353–356 (1993) doi: 10.1016/0013-4694(93)90049-2.
14. Kalauzi, A., Bojić, T., Vuckovic, A.: Modeling the Relationship Between Higuchi's Fractal Dimension and Fourier Spectra of Physiological Signals. *Med Biol Eng Comput*, 50(7), pp. 689–699 (2012) doi: 10.1007/s11517-012-0913-9. Epub 2012 May 17 PMID: 22588703.
15. Li, G., Xu, Y., Jiang, Y.: Mental Fatigue Has Great Impact on the Fractal Dimension of Brain Functional Network. *Neural Plasticity*, pp. 1–11 (2020) doi:10.1155/2020/8825547.
16. Cao, Z., Chuang, M., King, J. T.: Multi-Channel EEG Recordings During a Sustained-Attention Driving Task. *Scientific data*, 6(1), pp. 1–8 (2019) doi:10.1038/s41597-019-00274.
17. Pfurtscheller, G.: EEG Rhythms-Event-Related Desynchronization and Synchronization. *Rhythms in Physiological Systems*, pp. 289–296 (1991).
18. Pfurtscheller, G., Lopes da Silva, F.: Event-Related EEG/MEG Synchronization and Desynchronization: Basic Principles. *Clin Neurophysiol*, 110(11), pp. 1842–1857 (1999) doi: 10.1016/s1388-2457(99)00141-8. PMID: 10576479.
19. Mandelbrot, B., Wallis, J.: Robustness of The Rescaled Range R/S in the Measurement of Noncyclic Long Run Statistical Dependence. *Water Resources Research*, 5, pp. 967–988 (1969) doi: 10.1029/WR005i005p00967.
20. Delgado, O.Y., Delgado, J. R.: Hurst Exponent and Fractal Dimension Estimation of a Topographic Surface Through a Profiles Extraction. In: *UD y la geomática*, 5, pp. 84–91 (2011)
21. González-Rodríguez, D. G., Martínez-Peon, D., Angélica Ortiz Jiménez, X.: EEG-Based Drivers Mental Fatigue Detection Using ERD/ERS and Hurst Exponent. In: *Proceedings of the 6th International Conference on Medical and Health Informatics (ICMHI '22)*. Association for Computing Machinery, pp. 159–162 (2022) doi: 10.1145/3545729.3545763.

Emulación en FPGA de un memristor y su aplicación en un circuito caótico

Luis Jothan Gámez Palacios², Opeyemi Micheal Afolabi¹,
José Cruz Núñez Pérez¹

¹ Instituto Politécnico Nacional, IPN-CITEDI,
México

² Tecnológico Nacional de México,
Instituto Tecnológico de Tijuana,
México

luis.gamez@tectijuana.edu.mx, oafolabi@citedi.mx,
nunez@citedi.mx

Resumen. En este artículo se presenta el diseño y emulación de osciladores caóticos basados en memristores de orden entero y fraccional (IOMR y FOMR, respectivamente), implementados en el FPGA XC7A200T-2FBG676C mediante Vivado Design Suite. Los modelos matemáticos que describen el comportamiento de estos sistemas caóticos fueron validados mediante simulaciones numéricas, y posteriormente se diseñaron bloques HDL para compilar el modelo en VHDL. Para evaluar la respuesta y rendimiento del emulador, se analizan distintas configuraciones del oscilador caótico, variando el orden fraccional del sistema. El estudio de recursos y potencia en FPGA mostró que el oscilador caótico basado en IOMR demandó un 279% más potencia que el oscilador caótico basado en FOMR, cuya potencia nominal es de 1.42 W con un uso total de 15,891 recursos (entre LUTs, FFs, DSPs, IOs y BUFGs). La diferencia se atribuye a que el oscilador caótico basado en IOMR requiere 4.25 veces más recursos que el circuito oscilador basado en FOMR, debido a las técnicas de aproximación utilizadas para la solución de ambos sistemas caóticos.

Palabras Clave: FPGA, Memristor, Oscilador caótico, Simulink, VHDL, Vivado.

Memristor Emulation on FPGA and its Application in Chaotic System Circuits

Abstract. This paper presents the design and emulation of chaotic oscillators based on integer-order and fractional-order memristors (IOMR and FOMR, respectively), implemented on the XC7A200T-2FBG676C FPGA using the Vivado Design Suite. The mathematical models describing the behavior of these chaotic systems were validated through numerical simulations, and HDL blocks were subsequently designed to compile the model in VHDL. To evaluate the emulator's response and performance, different configurations of the chaotic oscillator were analyzed by varying the system's fractional order. FPGA resource

and power analysis showed that the IOMR-based chaotic oscillator required 279% more power than the FOMR-based oscillator, which had a nominal power consumption of 1.42 W and a total usage of 15,891 resources (including LUTs, FFs, DSPs, IOs, and BUFGs). The difference is attributed to the fact that the IOMR-based oscillator circuit demands 4.25 times more resources than the FOMR-based oscillator, due to the approximation techniques used to solve both chaotic systems.

Keywords. FPGA, Memristor, Chaotic oscillator, Simulink, VHDL, Vivado.

1. Introducción

En [1], Leon O. Chua analiza, desde la teoría de circuitos, los tres elementos básicos de los circuitos de dos terminales, definidos por relaciones entre dos de las cuatro variables fundamentales: voltaje (v), corriente (i), carga (q) y flujo magnético (ϕ). Tres de las seis relaciones posibles corresponden a la resistencia, el capacitor y el inductor. Chua argumenta que, desde una perspectiva matemática y simétrica, existe un vacío teórico en las relaciones restantes, lo que lo lleva a postular un cuarto elemento.

El concepto de dispositivo con memoria fue introducido por Chua en 1971 con la postulación teórica del memristor (MR), que presentaba una relación no lineal entre carga (q) y flujo magnético (ϕ), permitiéndole almacenar la carga que había fluido a través de él [2]. Aunque la teoría era sólida, no fue hasta 2008 que se logró crear físicamente el memristor en los laboratorios de HP bajo la dirección de Dmitri Strukov, confirmando las teorías de Chua casi 40 años después [3].

Para estudiar el comportamiento dinámico de circuitos con dispositivos con memoria, como los MRs, se han desarrollado modelos lineales por partes [4, 5] y macromodelos SPICE [6-8], útiles para la simulación de sistemas antes de su implementación física. En 2010, Pershin y Di Ventra [9] propusieron un emulador de MR basado en microcontroladores, lo que marcó un avance importante para el desarrollo de emuladores más sofisticados. Muthuswamy [10] propuso un emulador no lineal continuo suave basado en amplificadores operacionales y multiplicadores analógicos. En 2015, Yang *et al.* [11] desarrollaron un emulador de MR HP con características avanzadas como un amplio rango de memristancia, operación bimodal, largo período de no volatilidad y operación flotante. Zhao *et al.* [12] propusieron un emulador universal de memristores controlado por carga y voltaje, ofreciendo dos configuraciones principales, conectada a tierra o flotante, destacando por su versatilidad.

La integración de dispositivos con memoria en circuitos caóticos ha abierto un abanico de nuevas posibilidades de diseño y la implementación de sistemas con comportamientos dinámicos complejos introduciendo no linealidades y propiedades de memoria [13].

Este artículo está organizado de la siguiente manera, la sección 2 contiene teoría relevante para desarrollar los objetivos de este artículo, destacando los dispositivos con memoria, cálculo fraccional, métodos numéricos y sistemas caóticos. En la sección 3 se muestran las simulaciones numéricas en Matlab-Simulink de los osciladores caóticos basados en memristores y las métricas obtenidas del emulador implementado en FPGA mediante Vivado Design Suite. Por último, en la sección 4 se presentan las conclusiones

finales del artículo realizando comparaciones entre los resultados obtenidos con ambos programas.

2. Marco teórico

En esta sección se abordan los fundamentos teóricos necesarios para el desarrollo del presente trabajo. Se introduce el concepto de dispositivo con memoria, haciendo énfasis en el memristor, y se presentan los principios del cálculo fraccional y de sistemas caóticos.

2.1 Dispositivo de memoria

Los dispositivos con memoria, también conocidos como memelementos (ME), representan una clase de componente que posee la capacidad de almacenar su estado anterior debido a una relación no lineal entre ciertas variables, como voltaje, corriente, carga o flujo magnético. En [13] Di Ventra *et al.*, se establece que un ME controlado por m y descrito por un conjunto de k variables de estado x en un instante t se representa matemáticamente mediante la ecuación (1):

$$\begin{aligned} y(t) &= f(x, m, t) m(t), \\ \dot{x} &= g(x, m, t). \end{aligned} \quad (1)$$

donde $y(t)$ y $m(t)$ denotan la salida y la entrada del ME, respectivamente, f es la respuesta generaliza del sistema y g es una función vectorial continua en un espacio k dimensional.

Memristor. Di Ventra *et al.* [13] presentaron un modelo matemático para describir el comportamiento de un MR controlado por carga, en el cual la relación ideal que define su funcionamiento bajo la aplicación de un voltaje en sus terminales se expresa mediante la ecuación (2):

$$\begin{aligned} v_m(t) &= M(x, q, t) i(t), \\ \dot{x} &= g(x, q, t). \end{aligned} \quad (2)$$

donde M representa la memristancia que depende de la cantidad total de carga eléctrica y de la derivada de la función de flujo magnético con respecto a la carga acumulada en el tipo de MR considerando según la ecuación (3):

$$M(q) = \frac{d\phi(q)}{dq}. \quad (3)$$

El modelo matemático correspondiente al MR controlado por carga de la ecuación (2) se generaliza al orden arbitrario utilizando cálculo fraccional, obteniendo la expresión (4):

$$\begin{aligned} {}_0D_t^\alpha v_m(t) &= M(x, q, t) i(t), \\ \dot{x} &= g(x, q, t). \end{aligned} \quad (4)$$

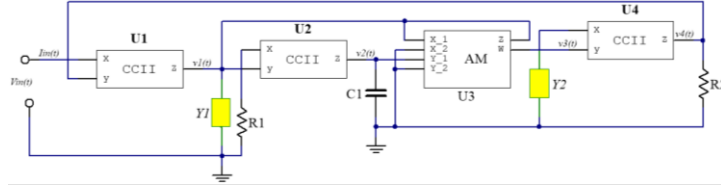


Fig. 1. Modelo universal de tipo conexión a tierra para memelementos de orden entero.

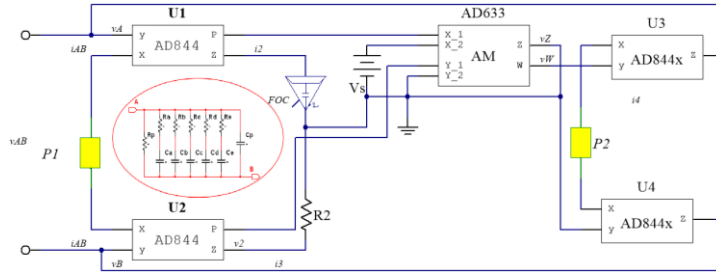


Fig. 2. Modelo universal de tipo conexión flotante para memelementos de orden fraccional.

Por otro lado, para el MR controlado por voltaje, el modelo matemático que define el comportamiento del dispositivo se expresa en la ecuación (5):

$$\begin{aligned} i_m(t) &= G(x, v, t) v(t), \\ \dot{x} &= g(x, v, t). \end{aligned} \quad (5)$$

donde G es el valor de la meminductancia del MR en el tiempo t . El dispositivo controlado por voltaje, al igual que el MR controlado por carga, se generaliza a valores arbitrarios tal como se describe en la ecuación (6):

$$\begin{aligned} {}_0 D_t^\alpha i_m(t) &= G(x, v, t) v(t), \\ \dot{x} &= g(x, v, t). \end{aligned} \quad (6)$$

Modelo universal de memelementos. En la obra de Zhao *et al.* [12], se propone un modelo universal de ME de tipo conexión a tierra de orden entero empleando transportadores de corriente de segunda generación de tipo positivo (CCII+), multiplicadores analógicos (AM), como se expone en la Fig. 1.

Las características terminales de las unidades CCII+ se describen en la ecuación (7), mientras que las características del AM se definen mediante la expresión (8):

$$v_x = v_y, \quad i_y = 0, \quad i_z = i_x, \quad (7)$$

$$v_w = \frac{(v_{x1} - v_{x2})(v_{y1} - v_{y2})}{10} + v_z. \quad (8)$$

Para un MR utilizando el modelo universal de la Fig. 1, los componentes Y_1 e Y_2 se sustituyen por los resistores R_0 y R_3 , respectivamente. El voltaje que atraviesa las terminales del MR de tipo conexión a tierra está dado por la expresión (9):

$$v_{in}(t) = -\frac{R_0 R_2}{R_3} i_{in}(t) + \frac{R_0^2 R_2}{10 R_1 R_3 C_1} q(t) i_{in}(t). \quad (9)$$

Por otro lado, el voltaje que atraviesa las terminales el MR de tipo conexión flotante se describe por la ecuación (10):

$$v_{in}(t) = \left(\frac{R_0 R_4}{R_5} - \frac{R_0 R_2}{R_3} \right) i_{in}(t) + \frac{R_0^2 R_2}{10 R_1 R_3 C_1} q(t) i_{in}(t). \quad (10)$$

En la obra de Li et al. [14], se propone un modelo universal de ME de tipo conexión flotante de orden fraccional empleando transportadores de corriente de segunda generación de tipo positivo y seguidores de tensión (CCII+), un multiplicador analógico (AM), un condensador de orden fraccionario (FOC), como se expone en la Fig. 2. Las propiedades de los AD833 y AD844 utilizados en el diseño se expresan mediante las ecuaciones (11) y (12), respectivamente:

$$v_x = v_y, \quad v_p = v_z, \quad i_y = 0, \quad i_z = i_x, \quad (11)$$

$$v_w = \frac{(v_{x1} - v_{x2})(v_{y1} - v_{y2})}{10} + v_z. \quad (12)$$

Para un MR utilizando el modelo universal de la Fig. 2, los componentes P_1 y P_2 se sustituyen por los resistores R_1 y R_3 , respectivamente. El corriente que atraviesa las terminales del MR de tipo conexión flotante está dado por la expresión (13):

$$i_{AB}(t) = \left(-\frac{R_2}{10 C_\alpha R_3 R_1} J^\alpha v_{AB}(t) - \frac{R_2 V_s}{10 R_1 R_3} \right) v_{AB}(t). \quad (13)$$

donde $v_{AB}(t)$ representa la entrada de voltaje. J^α es el operador de la integral fraccional de orden α , la cual se describe mediante la ecuación (14). V_s es la tensión de puerto del terminal x_2 del multiplicador AD633 y C_α es el valor del FOC:

$$J^\alpha v_{AB}(t) = \frac{1}{\Gamma(\alpha)} \int_0^t (t-\tau)^{\alpha-1} f(\tau) d\tau \quad (14)$$

2.2 Sistemas dinámicos

Los sistemas dinámicos son modelos matemáticos que describen la evolución de un sistema a lo largo del tiempo. Modelan las variables de estado en función del tiempo y las condiciones iniciales [15]. Los sistemas dinámicos de orden entero se describen mediante ecuaciones diferenciales ordinarias (ODE), como se describe en la ecuación (15):

$$\dot{x} = f(x, t, u), \quad (15)$$

donde $\dot{x}$, denota la derivada de x respecto al tiempo, f describe el comportamiento del sistema. Por otro lado, los sistemas dinámicos de orden fraccional se describen mediante ecuaciones diferenciales fraccionales (FDE), como se describe en la ecuación (16):

$$D_t^\alpha x(t) = f(x, t, u), \quad (16)$$

donde D_t^α es el operador de cálculo fraccional de orden α . El orden arbitrario puede ser un número real o racional y se define como $0 < \alpha \leq n$ y $f(x, t, u)$ representa la función que describe la dinámica del sistema.

2.3 Cálculo fraccional

El cálculo fraccional es una extensión natural de las matemáticas clásicas y se casi tan evidente tan pronto como se conocieron las ideas del cálculo básico de Leibniz, permitiendo explorar la viabilidad de aplicar potencias reales a los operadores integral y derivativo, en lugar de restringir dichos operadores a órdenes de naturaleza entera [15]. El origen del cálculo fraccional se remonta a 1695 con la notación $d^n f(x)/dx^n$ que introdujo Leibniz a L'Hospital para la derivada n -ésima de la función $f(x)$ para valores arbitrarios de n [15, 16]. Desde su introducción, se han estudiado muchos tipos de definiciones de derivadas e integrales de orden fraccional, por mencionar solo algunas, Laplace [17] obtuvo expresiones para ciertas derivadas fraccionales de casos particulares y Lacroix [16], quien este último proporcionó por primera vez una solución de derivada de orden m de una función potencia $f(x) = x^p$, tal como se desarrolla en la ecuación (17):

$$\frac{d^m}{dx^m}(x^p) = \frac{p!}{(p-m)!} x^{p-m}. \quad (17)$$

La aparición del operador factorial en la ecuación (17) sugiere la posibilidad de una generalización utilizando la función gamma, reescribiendo se obtiene la expresión (18):

$$\frac{d^m}{dx^m}(x^p) = \frac{\Gamma(p+1)}{\Gamma(p-m+1)} x^{p-m}, \quad (18)$$

donde m es el orden de la derivada la función x^p válida para exponentes p , mediante la relación $m \leq p$. La función gamma extiende la funcionalidad del número factorial k , denotado por $k!$ al dominio de números reales y complejos positivos [18]. La función gamma de un número complejo z se denota por $\Gamma(z)$, y es definida por la integral de la ecuación (19):

$$\Gamma(z) = \int_0^\infty e^{-t} t^{z-1} dt. \quad (19)$$

Definición de Riemann-Liouville. En 1832, Liouville estudió la derivada fraccional para $f(x) = e^{ax}$, pero su definición tenía restricciones de convergencia y generaba números complejos para ciertos casos. Optó por analizar primero la integral fraccional para luego definir la derivada fraccional. Paralelamente, Riemann obtuvo resultados similares mediante series de Taylor [15, 19]. Las investigaciones posteriores de

Bernhard Riemann y Joseph Liouville sobre los estudios que implican la tasa de cambio y las antiderivadas condujeron a la conocida derivada e integral de Riemann-Liouville de orden arbitrario, ecuaciones (20) y (21) respectivamente, para una función continua $f: [a, b] \rightarrow \mathbb{R}$, donde $m \in \mathbb{N}$, establecida por $(m-1) < \alpha \leq m$:

$${}^{\text{RL}}D_t^\alpha f(t) = \frac{1}{\Gamma(m-\alpha)} \left(\frac{d}{dt} \right)^m \int_a^t (t-k)^{m-\alpha-1} f(k) dk, \quad (20)$$

$${}^{\text{RL}}I_a^\alpha f(t) = \frac{1}{\Gamma(\alpha)} \int_a^t (t-k)^{\alpha-1} f(k) dk. \quad (21)$$

Definición de Caputo. Las definiciones de Riemann-Liouville son esenciales en el cálculo fraccional principalmente en el marco de un formalismo puramente matemático, pero presentan dificultades al modelar fenómenos reales [19]. En 1969, Michele Caputo propuso una nueva derivada fraccional, adecuada para problemas con valores iniciales o de frontera en sistemas físicos [18, 19]. La derivada fraccional de Caputo de orden $\alpha > 0$ de una función continua $f(t) \in L_1[a, b]$, con $a \geq 0$, es definida por la ecuación (22),

$${}^{\text{C}}D_t^\alpha f(t) = \begin{cases} \frac{1}{\Gamma(m-\alpha)} \int_a^t \frac{f^{(m)}(\tau)}{(t-\tau)^{\alpha+1-m}} d\tau, \\ \frac{d^m}{dt^m} f(t) \end{cases} \quad (22)$$

donde a y t son los límites del operador diferencial de orden fraccional, Γ denota la función gamma, el orden fraccional está definido por la $a=m$, y la expresión $(m-1) < \alpha \leq m$.

2.4 Métodos numéricos

Existen diversos métodos numéricos para aproximar la solución de sistemas de ecuaciones; sin embargo, cuando el problema a resolver es no lineal, estos métodos suelen implicar un alto costo computacional. Por esta razón, se emplean métodos paso a paso, como el método de Runge-Kutta y sus variantes, que permiten mejorar la eficiencia en la resolución de este tipo de problemas [20].

Runge – Kutta de orden entero. Fue desarrollado en el siglo XX por los matemáticos Carl Runge y Martin Kutta, basado en el método Euler para analizar ecuaciones diferenciales ordinarias. El método Runge – Kutta de cuarto orden (RK4) suele implementarse más a menudo debido a su alta tasa de convergencia para resolver problemas de la forma en la expresión (23), con condiciones iniciales en (24) definida en un dominio uniforme distante y un tamaño de paso $h = x_n - x_{n-1}$, donde $n=1, 2, \dots, N$ [18].

$$\frac{dy}{dx} = f(x, y), \quad (23)$$

$$y(x_0)=y_0. \quad (24)$$

El algoritmo de Runge – Kutta de cuarto orden está definido por la ecuación (25), donde k_n y h son coeficientes reales.

$$\begin{aligned} k_1 &= hf(x_n, y_n), \\ k_2 &= hf\left(x_n + \frac{h}{2}, y_n + \frac{k_1}{2}\right), \\ k_3 &= hf\left(x_n + \frac{h}{2}, y_n + \frac{k_2}{2}\right), \\ k_4 &= hf(x_n + h, y_n + k_3), \\ y_{n+1} &= y_n + \frac{(k_1 + 2k_2 + 2k_3 + k_4)}{6}. \end{aligned} \quad (25)$$

Runge – Kutta de orden fraccional. Considerando un problema con condición inicial definido por la derivada fraccional de Caputo en la ecuación (26), donde $0 < \alpha \leq 1$, $y(t) \in C[t_0, T]$, $f(t, y(t)) \in C[t_0, T]$ y t_0 es el punto base de la derivada fraccionaria [18]:

$$\begin{cases} {}^{C_0}D_t^\alpha y(t) = f(t, y(t)), & t \in [t_0, T], \\ y(t_0) = y_0. \end{cases} \quad (26)$$

El algoritmo Runge – Kutta explícito de 3 estados de orden fraccional (EFORK) se describe mediante la ecuación (27) con coeficientes y pesos determinados por la expresión (28):

$$\begin{aligned} k_1 &= h^\alpha f(t, y), \\ k_2 &= h^\alpha f(t + c_2 h, y + a_{21} k_1), \\ k_3 &= h^\alpha f(t + c_3 h, y + a_{31} k_1 + a_{32} k_2), \\ y_{n+1} &= y_n + w_1 k_1 + w_2 k_2 + w_3 k_3. \end{aligned} \quad (27)$$

$$\begin{aligned} a_{21} &= \frac{1}{2\Gamma(\alpha+1)^2}, \\ a_{31} &= \frac{\Gamma(1+\alpha)^2 \Gamma(2\alpha+1) + 2\Gamma(2\alpha+1)^2 \cdot \Gamma(3\alpha+1)}{4\Gamma(\alpha+1)^2 (2\Gamma(2\alpha+1)^2 \cdot \Gamma(3\alpha+1))}, \\ a_{32} &= -\frac{\Gamma(2\alpha+1)}{4(\Gamma(\alpha+1)^2 \cdot \Gamma(3\alpha+1))}, \\ c_2 &= \left(\frac{1}{2\Gamma(\alpha+1)}\right)^{1/\alpha}, \\ c_3 &= \left(\frac{1}{4\Gamma(\alpha+1)}\right)^{1/\alpha}, \\ w_1 &= \frac{8\Gamma(\alpha+1)^3 \Gamma(2\alpha+1)^2 - 6\Gamma(\alpha+1)^3 \Gamma(3\alpha+1) + \Gamma(2\alpha+1) \Gamma(3\alpha+1)}{\Gamma(\alpha+1) \Gamma(2\alpha+1) \Gamma(3\alpha+1)}, \\ w_2 &= \frac{2\Gamma(\alpha+1)^2 (4\Gamma(2\alpha+1)^2 - \Gamma(3\alpha+1))}{\Gamma(\alpha+1) \Gamma(3\alpha+1)}, \\ w_3 &= -\frac{8\Gamma(\alpha+1)^2 (2\Gamma(2\alpha+1)^2 - \Gamma(3\alpha+1))}{\Gamma(\alpha+1) \Gamma(3\alpha+1)}. \end{aligned} \quad (28)$$

2.5 Sistemas caóticos

Los sistemas caóticos son modelos dinámicos no lineales que presentan un comportamiento aparentemente impredecible, aunque sus fundamentos matemáticos

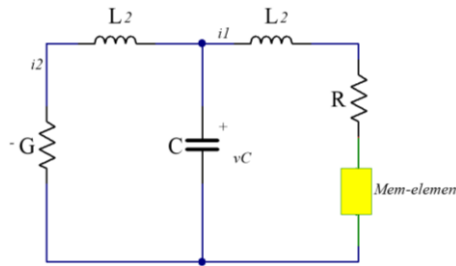


Fig. 3. Circuito oscilador caótico universal de memelementos de orden entero. sean completamente deterministas y se caracterizan por su sensibilidad a las condiciones iniciales [18].

Oscilador caótico universal de memelementos. En [12], Zhao *et al.* propusieron un circuito oscilador universal de ME de orden entero, expuesto en la Fig. 3. El circuito consiste de una conductancia negativa $-G$, los inductores L_1 y L_2 , un resistor R , un capacitor C , y un modelo de ME caracterizado por su voltaje $v_m(q)$.

Aplicando ley de Kirchhoff, el sistema dinámico del circuito oscilador para un ME se deriva en la ecuación (29),

$$\begin{cases} L_1 \frac{di_1}{dt} = v_c - Ri_1 - v_m(t), \\ C \frac{dv_c}{dt} = i_2 - i_1, \\ L_2 \frac{di_2}{dt} = \frac{i_2}{G} - v_c, \\ \frac{dq}{dt} = i_1. \end{cases} \quad (29)$$

Cuando el ME se sustituye por un IOMR en el circuito de la Fig. 3, y empleando la ecuación (9), se deriva la ecuación (30) donde i_1 , i_2 , q , y v_c son variables de estado en el sistema dinámico.

$$\begin{cases} L_1 \frac{di_1}{dt} = v_c - Ri_1 - \left(-\frac{R_0 R_2}{R_3} + \frac{R_0^2 R_2}{10 R_1 R_3 C_1} q \right), \\ C \frac{dv_c}{dt} = i_2 - i_1, \\ L_2 \frac{di_2}{dt} = \frac{i_2}{G} - v_c, \\ \frac{dq}{dt} = i_1, \end{cases} \quad (30)$$

Realizando un escalamiento en el dominio del tiempo del sistema de la ecuación (30), la representación se simplifica al sistema dinámico de la ecuación (31) utilizando $t = L_2 d\tau$, $x = i_1$, $y = v_c$, $z = i_2$, $w = q/L_2$, $a = L_1/L_2$, $b = L_2/C$, $c = R$, $k = 1/G$, $m = -R_0 R_2/R_3$ y $n = R_0^2 R_2/R_1 R_3 C_1$, donde x, y, z, w son las variables de estado y a, b, c, k, m, n , son valores constantes reales.

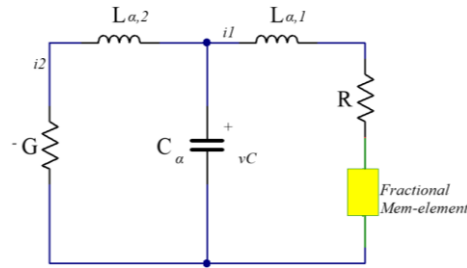


Fig. 4. Circuito oscilador caótico universal de memelementos de orden fraccional.

$$\begin{cases} \dot{x} = a(y - cx - (m + nw)x), \\ \dot{y} = b(z - x), \\ \dot{z} = kz - y, \\ \dot{w} = x, \end{cases} \quad (31)$$

Al considerar el FOMR, la topología del circuito de la Fig. 3 se modifica por elementos de orden arbitrario. Zhao *et al.* proponen en [12] un circuito oscilador universal de ME que considera el orden arbitrario, expuesto en la Fig. 4. El circuito consiste de un capacitor de orden fraccional C_α , los inductores $L_{\alpha,1}$ y $L_{\alpha,2}$, un resistor R , una conductancia negativa $-G$ y un modelo de ME de orden fraccional caracterizado por su voltaje $v_m(q)$.

La dinámica del circuito oscilador caótico se modela matemáticamente como un sistema fraccional conmensurado utilizando leyes de Kirchhoff como se describe la ecuación (32),

$$\begin{cases} L_{\alpha,1} {}^c D_{t_0}^\alpha i_1 = v_C - Ri_1 - v_m(t), \\ C_\alpha {}^c D_{t_0}^\alpha v_C = i_2 - i_1, \\ L_{\alpha,2} {}^c D_{t_0}^\alpha i_2 = \frac{i_2}{G} - v_C, \\ {}^c D_{t_0}^\alpha q = i_1, \end{cases} \quad (32)$$

Sustituyendo el ME por un FOMR y realizando un escalamiento en el dominio del tiempo, como en las ecuaciones (30) y (31), se obtiene el sistema dinámico con la incorporación del orden arbitrario en la ecuación (33),

$$\begin{cases} {}^c D_{t_0}^\alpha x = a(y - cx - (m + nw)x), \\ {}^c D_{t_0}^\alpha y = b(z - x), \\ {}^c D_{t_0}^\alpha z = kz - y, \\ {}^c D_{t_0}^\alpha w = x, \end{cases} \quad (33)$$

donde x, y, z, w son las variables de estado y a, b, c, k, m, n , son valores constantes reales.

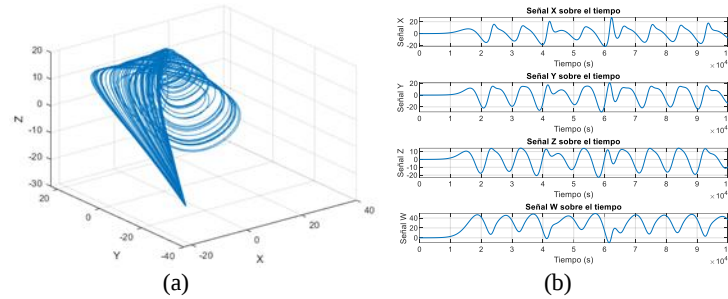


Fig. 5. Respuesta del emulador para un oscilador caótico basado en IOMR: a) gráfico 3D de la relación $x, y, y z$, b) variables sobre el tiempo.

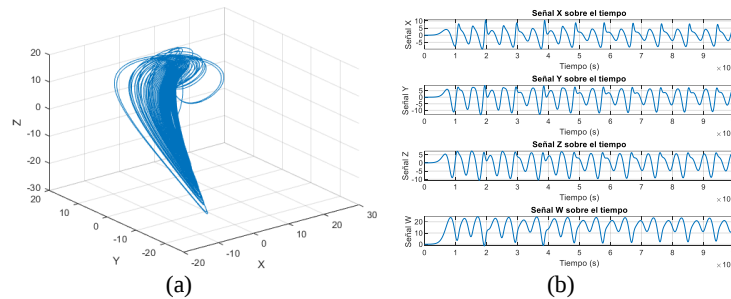


Fig. 6. Respuesta del emulador para un oscilador caótico basado en FOMR cuando $\alpha=0.9$: a) gráfico 3D de la relación $x, y, y z$, b) variables sobre el tiempo.

3. Diseño digital de los osciladores caóticos basados en memristores

En este trabajo se diseñó mediante bloques HDL los osciladores caóticos basados en memristores de orden entero y fraccional en la plataforma de Matlab-Simulink, permitiendo compilar el diseño en VHDL para implementarlo en hardware mediante Vivado Design Suite.

3.1 Oscilador caótico basado en IOMR

El oscilador caótico basado en IOMR descrito en la ecuación (31), se define bajo los siguientes parámetros: $a=2$, $b=1$, $c=0.2$, $k=0.92$, $m=-0.002$ y $n=0.04$. El sistema se inicializa con los estados $x(0)=y(0)=z(0)=w(0)=0.01$. Para aproximar la solución al sistema de ecuaciones, se emplea RK4, con un tamaño de paso $h=0.001$. En la Fig. 5a se expone la relación de $x, y, y z$ que rigen la dinámica del oscilador caótico en 3D, mientras que en la Fig. 5b se muestra el comportamiento de las variables sobre el tiempo obtenidas por el emulador.

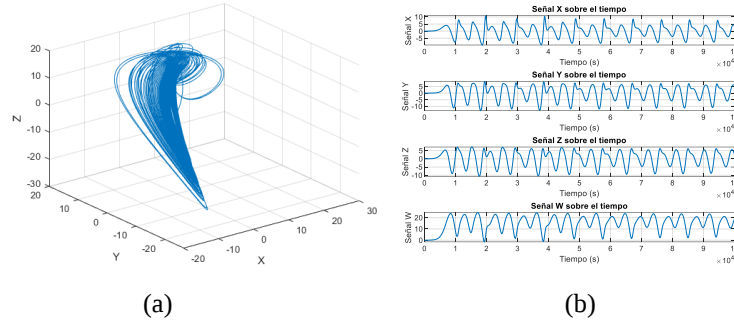


Fig. 7. Respuesta del emulador para un oscilador caótico basado en FOMR cuando $\alpha=0.95$: a) gráfico 3D de la relación x , y , y z , b) variables sobre el tiempo.

Tabla 1. Consumo de potencia del oscilador caótico basado en IOMR emulado en FPGA.

	Distribución de potencia	Oscilador caótico basado en IOMR
Dinamic	Clocks	0.006 W
	Signals	1.870 W
	Logic	1.259 W
	DSP	0.593 W
	I/O	0.090 W
Device static	---	0.143 W
	Total	3.960 W

3.2 Oscilador caótico basado en FOMR

El oscilador caótico basado en FOMR descrito en la ecuación (33), se define bajo los siguientes parámetros: $a=2$, $b=1$, $c=0.2$, $k=0.92$, $m=-0.002$ y $n=0.04$ para $\alpha=0.9$ y $\alpha=0.95$. El sistema se inicializa con los estados $x(0)=y(0)=z(0)=w(0)=0.01$. Para aproximar la solución al sistema de ecuaciones, se emplea EFORK, con un tamaño de paso $h=0.001$. En la Fig. 6a se expone la relación de x , y , y z que rigen la dinámica del oscilador caótico en 3D para $\alpha=0.9$, mientras que en la Fig. 6b se muestra el comportamiento de las variables sobre el tiempo obtenidas por el emulador.

En la Fig. 7a se expone la relación de x , y , y z que rigen la dinámica del oscilador caótico en 3D para $\alpha=0.95$, mientras que en la Fig. 7b se muestra el comportamiento de las variables sobre el tiempo obtenidas por el emulador.

4. Resultados

En esta sección se presentan los resultados obtenidos de la implementación del circuito oscilador caótico basado en memristores, tanto para el caso de orden entero como para el orden fraccional, utilizando una FPGA como plataforma de prueba. Se

Tabla 2. Consumo de potencia del oscilador caótico basado en FOMR cuando $\alpha=0.9$ emulado en FPGA.

	Distribución de potencia	Oscilador caótico basado en IOMR
Dinamic	Clocks	0.005 W
	Signals	0.486 W
	Logic	0.276 W
	DSP	0.422 W
	I/O	0.090 W
Device static	---	0.134 W
	Total	1.413 W

Tabla 3. Consumo de potencia del oscilador caótico basado en FOMR cuando $\alpha=0.95$ emulado en FPGA.

	Distribución de potencia	Oscilador caótico basado en IOMR
Dinamic	Clocks	0.005 W
	Signals	0.490 W
	Logic	0.277 W
	DSP	0.424 W
	I/O	0.090 W
Device static	---	0.134 W
	Total	1.420

comparan ambas implementaciones en términos de métricas de desempeño, considerando la cantidad de recursos utilizados y la potencia requerida por el FPGA.

4.1 Oscilador caótico basado en IOMR

Tras la implementación exitosa en FPGA del emulador, se analizan las métricas de este último. En la Tabla 1 se expone la distribución del consumo de potencia que demanda el emulador, utilizando un total de 3.96 W. El emulador emplea 67,553 recursos del FPGA, repartidos de la siguiente manera: 66,497 LUTs (49.7%), 168 FFs (0.06%), 717 DSPs (96.8%), 170 IO (42.5%) y 1 BUFG (3.13%).

4.2 Oscilador caótico basado en FOMR

En la Tabla 2 se expone la distribución del consumo de potencia que demanda el emulador, utilizando un total de 3.96 W. El emulador emplea 15,891 recursos del FPGA, repartidos de la siguiente manera: 14,932 LUTs (11.6%), 287 FFs (0.001%), 501 DSPs (67.7%), 170 IO (42.5%) y 1 BUFG (3.13%).

En la Tabla 3 se expone la distribución del consumo de potencia que demanda el emulador, utilizando un total de 1.42 W. Los recursos utilizados por el emulador para $\alpha=0.95$ coinciden con los utilizados por el emulador para $\alpha=0.9$.

5. Conclusiones

En este trabajo se compararon osciladores caóticos basados en IOMR y FOMR con órdenes arbitrarios $\alpha=0.9$ y $\alpha=0.95$. La emulación en FPGA permitió analizar el uso de recursos, consumo de potencia y eficiencia computacional. Se encontró que el oscilador basado en IOMR requiere más recursos que el de FOMR, debido al método de aproximación empleado. Aunque el método RK4 aumenta las iteraciones y el uso de recursos, el método EFORK de tres etapas optimizó parcialmente la demanda computacional. El oscilador basado en IOMR consume un 279% más de potencia que el de FOMR. En términos de lógica programable, el oscilador IOMR ocupa el 49.7% de las *LUTs* y el 96.8% de los *DSPs*, mientras que el FOMR solo usa el 11.16% de las *LUTs* y el 67.7% de los *DSPs*, evidenciando una mayor eficiencia en la implementación del FOMR en FPGA.

Agradecimientos. Los autores agradecen al Instituto Politécnico Nacional por el apoyo recibido a través del proyecto SIP 20251203.

Referencias

1. Chua, L.: Memristor-The Missing Circuit Element. *IEEE Transactions on Circuit Theory*, 18(5), pp. 507–519 (1971) doi: 10.1109/TCT.1971.1083337.
2. Khalil, N. A., Fouda, M. E., Said, L. A.: A Universal Fractional-Order Memelement Emulation Circuit. In: 2019 Novel Intelligent and Leading Emerging Sciences Conference (NILES), pp. 67–70 (2019) doi: 10.1109/NILES.2019.8909307.
3. Renping; W., Chunhua, W.: A New Simple Chaotic Circuit Based on Memristor. *International Journal of Bifurcation and Chaos*, 26(9) pp. 1650145 (2016) doi:10.1142/S0218127416501455.
4. Itoh M., Chua, L. O.: Memristor Oscillators. *International Journal of Bifurcation and Chaos*, 18(11), pp. 3183–3206 (2008).
5. Muthuswamy B., Kokate, P. P.: Memristor-based Chaotic Circuits. *IETE Technical Review*, 26, pp. 415–426 (2009)
6. Benderli S., Wey, T. A.: On SPICE Macromodelling of TiO₂ Memristor. *Electronics Letters*, 45, pp. 377–378 (2009) doi: 10.1049/el.2009.3511.
7. Rak A., Cserey, G.: Macromodeling of The Memristor in SPICE. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, 29(4), pp. 632–636 (2010) doi: 10.1109/TCAD.2010.2042900.
8. Batas D., Fiedler, H.: A Memristor SPICE Implementation and A New Approach for Magnetic Flux-Controlled Memristor Modeling. *IEEE Transactions on Nanotechnology*, 10(2), pp. 250–255 (2011)
9. Pershin Y. V., Di Ventra, M.: Practical Approach to Programmable Analog Circuits with Memristors. *IEEE Transactions on Circuits and Systems I: Regular Papers*, 57(8), pp. 1857–1864 (2010)
10. Muthuswamy B., Chua, L. O.: Simplest Chaotic Circuit. *International Journal of Bifurcation and Chaos*, 20(5), pp. 1567–1580 (2010)
11. Yang, C. J., Choi, H., Park, S.: A memristor Emulator as A Replacement of A Real Memristor. *Semiconductor Science and Technology*, 30, pp. 1–9 (2015)
12. Zhao, Q., Wang, C., Zhang, X.: A Universal Emulator for Memristor, Memcapacitor, and Meminductor and Its Chaotic Circuit. *Chaos: An Interdisciplinary Journal of Nonlinear Science*, 29(1), pp. 013–141 doi:10.1063/1.5081076.

13. Di Ventra, M., Pershin, Y. V., Chua, L. O.: Circuit Elements with Memory: Memristors, Memcapacitors, and Meminductors. *Proceedings of the IEEE*, 97(10), pp. 1717–1724 (2009) doi: 10.1109/JPROC.2009.2021077.
14. Li, Y., Xie, L., Zheng, C.: Modeling and Hardware Implementation of Universal Interface-Based Floating Fractional-Order Memelements. *Chaos: An Interdisciplinary Journal of Nonlinear Science*, 33(1) (2023) doi: 10.1063/5.0124793.
15. Strogatz, S. H.: *Nonlinear Dynamics and Chaos: With Applications to Physics, Biology, Chemistry, and Engineering*, 2nd ed. Boca Raton, FL, USA: CRC Press (2015)
16. Podlubny, I.: *Generalized Viscoelastic Models: Their Fractional Equations with Solutions*, vol. 198. New York, NY, USA: Academic Press (1998)
17. Ross, B.: The Development of Fractional Calculus 1695–1900, *Historia Mathematica*, 4(1), pp. 75–89 (1977) doi: 10.1016/0315-0860(77)90039-8.
18. Arfken, G. B., Weber, H. J., Harris, F. E.: *Mathematical Methods for Physicists*, 5th ed. San Diego, CA, USA: Academic Press (2001)
19. Oldham, K. B., Spanier, J.: *The Fractional Calculus*, vol. 17. New York, NY, USA: Dover, 1st ed. (1974)
20. Ghoreishi, F., Ghaffari, R., Saad, N.: Fractional Order Runge–Kutta Methods. *Fractal Fract.*, 7(3), pp. 245 (2023) doi: 10.3390/fractalfract7030245.

Electronic edition
Available online: <http://www.rcs.cic.ipn.mx>



<http://rsc.cic.ipn.mx>



Centro de Investigación
en Computación