

Special Issue Reprint

Fuzzy Decision Making and Soft Computing Applications

Future Perspectives

Edited by
Sachi Nandan Mohanty

mdpi.com/journal/mathematics

Fuzzy Decision Making and Soft Computing Applications: Future Perspectives

Fuzzy Decision Making and Soft Computing Applications: Future Perspectives

Guest Editor

Sachi Nandan Mohanty



Basel • Beijing • Wuhan • Barcelona • Belgrade • Novi Sad • Cluj • Manchester

Guest Editor

Sachi Nandan Mohanty
School of Computer Science
& Engineering
VIT-AP University
Amaravati
India

Editorial Office

MDPI AG
Grosspeteranlage 5
4052 Basel, Switzerland

This is a reprint of the Special Issue, published open access by the journal *Mathematics* (ISSN 2227-7390), freely accessible at: https://www.mdpi.com/si/mathematics/Fuzzy_Decis_Mak_Soft_Comput.

For citation purposes, cite each article independently as indicated on the article page online and as indicated below:

Lastname, A.A.; Lastname, B.B. Article Title. <i>Journal Name</i> Year , Volume Number, Page Range.
--

ISBN 978-3-7258-4657-3 (Hbk)

ISBN 978-3-7258-4658-0 (PDF)

<https://doi.org/10.3390/books978-3-7258-4658-0>

© 2025 by the authors. Articles in this book are Open Access and distributed under the Creative Commons Attribution (CC BY) license. The book as a whole is distributed by MDPI under the terms and conditions of the Creative Commons Attribution-NonCommercial-NoDerivs (CC BY-NC-ND) license (<https://creativecommons.org/licenses/by-nc-nd/4.0/>).

Contents

About the Editor	vii
Preface	ix
Leticia Cervantes, Camilo Caraveo and Oscar Castillo	
Performance Comparison between Type-1 and Type-2 Fuzzy Logic Control Applied to Insulin Pump Injection in Real Time for Patients with Diabetes	
Reprinted from: <i>Mathematics</i> 2023 , <i>11</i> , 730, https://doi.org/10.3390/math11030730	1
Arihant Tanwar, Wajdi Alghamdi, Mohammad D. Alahmadi, Harpreet Singh and Prashant Singh Rana	
A Fuzzy-Based Fast Feature Selection Using Divide and Conquer Technique in Huge Dimension Dataset	
Reprinted from: <i>Mathematics</i> 2023 , <i>11</i> , 920, https://doi.org/10.3390/math11040920	14
Dibyalekha Nayak, Kananbala Ray, Tejaswini Kar and Sachi Nandan Mohanty	
Fuzzy Rule Based Adaptive Block Compressive Sensing for WSN Application	
Reprinted from: <i>Mathematics</i> 2023 , <i>11</i> , 1660, https://doi.org/10.3390/math11071660	29
Manuel Casal-Guisande, Jorge Cerqueiro-Pequeño, José-Benito Bouza-Rodríguez and Alberto Comesaña-Campos	
Integration of the Wang & Mendel Algorithm into the Application of Fuzzy Expert Systems to Intelligent Clinical Decision Support Systems	
Reprinted from: <i>Mathematics</i> 2023 , <i>11</i> , 2469, https://doi.org/10.3390/math11112469	50
Hector Carreon-Ortiz, Fevrier Valdez and Oscar Castillo	
Comparative Study of Type-1 and Interval Type-2 Fuzzy Logic Systems in Parameter Adaptation for the Fuzzy Discrete Mycorrhiza Optimization Algorithm	
Reprinted from: <i>Mathematics</i> 2023 , <i>11</i> , 2501, https://doi.org/10.3390/math11112501	83
Shweta Singhal, Nishtha Jatana, Kavita Sheoran, Geetika Dhand, Shaily Malik, Reena Gupta, et al.	
Multi-Objective Fault-Coverage Based Regression Test Selection and Prioritization Using Enhanced ACO-TCSP	
Reprinted from: <i>Mathematics</i> 2023 , <i>11</i> , 2983, https://doi.org/10.3390/math11132983	121
Ronald Díaz Cazañas, Daynier Rolando Delgado Sobrino, Estrella María De La Paz Martínez, Jana Petru and Carlos Daniel Díaz Tejada	
Proposal of a Framework for Evaluating the Importance of Production and Maintenance Integration Supported by the Use of Ordinal Linguistic Fuzzy Modeling	
Reprinted from: <i>Mathematics</i> 2024 , <i>12</i> , 338, https://doi.org/10.3390/math12020338	142
Ahmet Topal, Nilgun Guler Bayazit and Yasemen Ucan	
A Method to Handle the Missing Values in Multi-Criteria Sorting Problems Based on Dominance Rough Sets	
Reprinted from: <i>Mathematics</i> 2024 , <i>12</i> , 2944, https://doi.org/10.3390/math12182944	167
Ch Sree Kumar, Aayushman Bhaba Padhy, Akhilendra Pratap Singh and K. Hemant Kumar Reddy	
A Dynamic Trading Approach Based on Walrasian Equilibrium in a Blockchain-Based NFT Framework for Sustainable Waste Management	
Reprinted from: <i>Mathematics</i> 2025 , <i>13</i> , 521, https://doi.org/10.3390/math13030521	185

About the Editor

Sachi Nandan Mohanty

Sachi Nandan Mohanty has been recognized in the Top 2% World Scientists Ranking by Stanford University and Elsevier for 2022, 2023, and 2024. He earned his Ph.D. from the Indian Institute of Technology (IIT) Kharagpur in 2015 with an MHRD scholarship from the Government of India, and completed his Postdoctoral research at IIT Kanpur in 2019. Dr. Mohanty has authored and edited 42 books published by IEEE-Wiley, Springer, Wiley, CRC Press, NOVA, and DeGruyter. His research expertise spans data mining, big data analysis, cognitive science, fuzzy decision-making, brain-computer interfaces, cognition, and computational intelligence. He has published 241 research papers in reputed international journals and has guided nine Ph.D. scholars and 23 postgraduate students. During his Ph.D., he received four Best Paper Awards. His thesis earned the Best Thesis Award (First Prize) from the Computer Society of India in 2015. He has received international travel support four times from the Government of India's Department of Science and Technology (DST-SERB) to present his research papers and deliver keynote speeches. Dr. Mohanty has been honored with numerous awards and fellowships, including the Prof. Ganesh Mishra Memorial Award and the 19th Prof. Bhubaneswar Behera Lecture Award. He is a Fellow of the Institution of Engineers (IEI) and IETE, an Ambassador of the European Alliance for Innovation (EAI), and a Senior Member of the IEEE Computer Society (Hyderabad Chapter). Additionally, he serves as the General Chair for the international conferences ICISML and AIHC and as the Editor-in-Chief of *EAI Transactions on Intelligent Systems and Machine Learning Applications*, *Journal of Sustainable Development Innovations*, and *International Journal on Smart & Sustainable Intelligent Computing*. His academic contributions have taken him to prestigious conferences across the U.S., Europe, the Middle East, and Asia.

Preface

The increasing complexity of modern systems and the uncertainty inherent in real-world data have driven researchers and practitioners to explore advanced decision-making tools grounded in soft computing and fuzzy logic. This volume presents a curated collection of innovative research contributions that reflect the latest advancements in fuzzy decision-making, optimization, and intelligent systems across diverse application domains.

The article titled "A Dynamic Trading Approach Based on Walrasian Equilibrium in a Blockchain-Based NFT Framework for Sustainable Waste Management" introduces a novel integration of economic theory with blockchain technology to foster sustainable practices. In "A Method to Handle the Missing Values in Multi-Criteria Sorting Problems Based on Dominance Rough Sets," the focus is placed on enhancing the robustness of decision-making frameworks in the presence of incomplete data.

Addressing industrial challenges, "Proposal of a Framework for Evaluating the Importance of Production and Maintenance Integration Supported by the Use of Ordinal Linguistic Fuzzy Modeling" explores the synergy between operations and maintenance through a fuzzy linguistic approach. The article "Multi-Objective Fault-Coverage Based Regression Test Selection and Prioritization Using Enhanced ACO.TCSP" leverages nature-inspired algorithms for the improved effectiveness of software testing.

In the realm of computational intelligence, "Comparative Study of Type-1 and Interval Type-2 Fuzzy Logic Systems in Parameter Adaptation for the Fuzzy Discrete Mycorrhiza Optimization Algorithm" offers insights into fuzzy system adaptability. Similarly, "Integration of the Wang & Mendel Algorithm into the Application of Fuzzy Expert Systems to Intelligent Clinical Decision Support Systems" showcases the use of fuzzy rule-based inference in healthcare.

Sensor networks and data compression are addressed in "Fuzzy Rule Based Adaptive Block Compressive Sensing for WSN Application," proposing a smart solution for efficient data transmission. The article "A Fuzzy-Based Fast Feature Selection Using Divide and Conquer Technique in Huge Dimension Dataset" contributes to the field of data mining by enhancing feature selection methods for high-dimensional datasets.

Finally, "Performance Comparison between Type-1 and Type-2 Fuzzy Logic Control Applied to Insulin Pump Injection in Real Time for Patients with Diabetes" presents a critical evaluation of fuzzy control strategies in biomedical applications, underlining the real-time impact of intelligent systems on human health.

Together, these contributions illustrate the versatility and power of fuzzy and soft computing approaches in solving multifaceted challenges. We hope this volume will serve as a valuable reference for researchers, practitioners, and students engaged in the pursuit of intelligent and sustainable technological solutions.

Sachi Nandan Mohanty

Guest Editor

Article

Performance Comparison between Type-1 and Type-2 Fuzzy Logic Control Applied to Insulin Pump Injection in Real Time for Patients with Diabetes

Leticia Cervantes ¹, Camilo Caraveo ^{1,*} and Oscar Castillo ²¹ Autonomous University of Baja California, Tijuana 21500, Mexico² Tijuana Institute of Technology, Tijuana 22414, Mexico

* Correspondence: camilo.caraveo@uabc.edu.mx

Abstract: Nowadays, type 1 diabetes is unfortunately one of the most common diseases, and people tend to develop it due to external factors or by hereditary factors. If is not treated, this disease can generate serious consequences to people's health, such as heart disease, neuropathy, pregnancy complications, eye damage, etc. Stress can also affect the condition of patients with diabetes, and our motivation in this work is to help manage the health of people with type 1 diabetes. The contribution of this paper is in presenting the implementation of type-1 and type-2 fuzzy controllers to control the insulin dose to be applied in people with type 1 diabetes in real time and in stressful situations. First, a diagram for the insulin control is presented; second, type-1 and type-2 fuzzy controllers are designed and tested on the insulin pump in real time over a 24 h period covering one day; then, a comparative analysis of the performance of these two controllers using a statistical test is presented with the aim of maintaining a stable health condition of people through an optimal insulin supply. In the model for the insulin control, perturbations (noise/stress levels) were added to find if our proposed fuzzy controller has good insulin control in situations that could generate disturbances in the patient, and the results found were significant; in most of the tests carried out, the type-2 controller proved to be more stable and efficient; more information can be found in the discussion section.

Keywords: fuzzy logic; artificial pancreas; intelligent control; continuous insulin pump

MSC: 03-08

1. Introduction

At present, technology has evolved a great deal, and several applications have been made using intelligent techniques, for example, in the optimization of processes and reduction of times applied in industry and also in some process to improve the quality of life for people. For example, mechatronic devices can be a physical part of people, for example, an arm or leg; smaller devices can measure heart rate and notify an emergency contact if the person's health is in danger or even call automatically to the ambulance, and there are also continuous pressure monitoring devices that provide more control over a person's health, especially older people. Many processes that were previously manual today have been automated for better monitoring and control for people, such as the artificial pancreas, which is an automated process of a human organ that helps regulate a specific function to improve people's quality of life. We know that the most important thing for human beings today is health because without it, we have nothing. One can enjoy money, comforts, friends, and trips, but if one does not have health, one cannot enjoy any of these things. There are many factors and different diseases that attack human life, and some of them can be controlled with medication, while some of them cannot. Hereditary diseases or diseases caused by an external factor also have a great impact on people's living

conditions and state of mind. One of the most common diseases is diabetes, and there are different variants of this disease [1].

In this research, the focus is specifically for type 1 diabetes. It was mentioned that technology is used for various benefits, including people's health, and on this occasion, for the control of diabetes, there is an insulin pump. These devices are used to avoid injections and make it easier to administer insulin doses. A good controller is needed to have a good control of the daily insulin dose in real time for patients, and this is where this research focuses: on developing type-1 and type-2 fuzzy systems for controlling the insulin dose in real time.

It is important to mention that in the literature, there are few works related to our proposal, and the most relevant are mentioned below. In [2], the authors of this work developed an embedded system using an FPGA microcontroller for insulin injection; however they did not integrate intelligent techniques into their work. Furthermore, in [3], a case study of insulin injection pumps was presented to analyze the behavior of users of different ages and stages of diabetes in a controlled space. In [4], a comparison was performed of two hybrid closed-loop systems in adolescents and young adults with type 1 diabetes (FLAIR). Here, a multicenter, randomized, crossover trial is presented, and the authors of this paper carry out a comparative study between different countries with the aim of testing whether the new automatic methods are more efficient than traditional manual methods.

This paper is organized as follows: Section 2 presents the basic concepts to better understand the application. The description of the problem and the methodology are presented in Section 3. Results of type-1 and type-2 fuzzy systems and their comparison are presented in Section 4, and finally, in Section 5, the conclusions are presented.

2. Background and Basic Concepts

Some basic concepts such as type 1 diabetes, type-1 and type-2 fuzzy systems, artificial pancreas, intelligent control, and continuous insulin pump are explained in this section to understand better the application and the contribution.

2.1. Type 1 Diabetes

Type 1 diabetes disease is considered the second most chronic disease and most frequently presents in childhood [5]. When we talk about type 1 diabetes, we are referring to a condition where the pancreas does not produce insulin, and this means that there is an excess of glucose that remains in the blood [6], for which it is important to take into account the desired glucose levels [7]. If these problems persist, and no attention is paid, it can generate several serious problems, for example, in the heart [8], eyes, kidneys, etc. In these times of the COVID-19 pandemic, some patients who had certain diseases [9] including diabetes showed greater complications when infected and were more vulnerable [10,11].

Constant, real-time monitoring is essential to make medical decisions; if a real-time monitor is added for each patient, with personalized insulin supply control for each person, better results can be obtained. Today, there are various technologies that have been added to treat and control type 1 diabetes. Ordinarily, intramuscular insulin supply was common, which is one of the best known approaches, but there are also glucose monitoring devices. Glucose monitoring allows people to know if their glycemic goals are being met. Previously, the glucose level was measured in the urine using a copper reagent, but today, there are portable devices that can measure it based on a drop of blood or by using a continuous monitoring device. In the search for alternatives for patients with diabetes, a predictive suspension pumps with low glucose level (PLGS) can be found; these devices that use this technology interrupt the administration of insulin when it is detected so that the glucose value of the sensor reaches or decreases to the glucose limit within a period of time [12].

There are several kinds of devices that are coming into the market with the objective of facilitating glucose control for people with diabetes [13]. Moreover, it is important to mention that different technological tools, computing techniques, and mathematical

models are used to research and develop new advances for people with type 1 diabetes. For example, in [14], mathematical models are presented to help to understand and predict the dynamics within the glucose regulation system, and this part is very important because based on these models' possible behaviors, methods can be developed and contribute to the health of people and to other work, where they also use techniques to predict glucose concentration using image processing and machine learning [15]. There is also what today is called an artificial pancreas, which is glucose-sensitive automated insulin administration and works using computer techniques such as PID control, i.e., combining glucose control with insulin administration; in short, it is a system that provides automatic control of blood glucose levels and assists in the administration of insulin. Several studies have been carried out, and it has been seen that insulin pumps help to improve the quality of people's lives [16] since they provide better control in daily social activities as well as the use of the different interfaces developed as in [17], which presents two ways to implement an MPC (model predictive control) for a smartphone-based artificial pancreas system. There are also investigations where the authors [18] performed a system based on derivatives of fractional order with the objective of regulating the supply for an insulin pump with the objective of regulating diabetes. Therefore, though we can continue giving different examples, it is easy to see that researchers want to reach the same goal, which is to obtain support for people with type 1 diabetes.

Diabetes and Stress

There are many myths that say that a scare can cause diabetes, but there is no established foundation; however, for people who have this tendency, increased stress can accelerate the beginning of this disease, and a person who is constantly subjected to stress can develop health problems. Stress is characterized by violent nervous tension that is maintained in addition to being accompanied by a significant degree of anxiety, and it can be said that there are different types of stress such as psychological, social, economic, physiological, and psychosocial [19,20]. When a person is under stress levels, they can have symptoms such as headache, chest and back and neck pain, cold sweats, insomnia, etc., in some systems such as the metabolic, cardiovascular, and gastrointestinal systems, among others. Only in Latin America is work stress a psychosocial factor, and it is considered an epidemic of modern working life. When some people are under stress, they are susceptible to consuming an excess of calories, which can cause weight gain and variations in their glucose levels, which is especially concerning in people with diabetes. In addition, the emotional stress generated by living with a person with diabetes can negatively affect treatment adherence, quality of life, and disease control [21].

Stress can cause an increase in blood glucose levels. The increase in stress at chronic levels will depend on the vulnerability of each person, their ability to protect and respond to the body, and their self-esteem as well as the social support they have [22]. Each of these factors is very important because each can result in a stress with high levels, which it can generate a significant change in glucose, and if the control is not the adequate at that time, it could generate significant changes in people's health. Different levels of stress have different effects on each person. For example, some hardly eat, while others eat an excessive amount, and some turn to junk food or try to over-exercise when they feel stressed to try to release it. However, there is also oxidative stress, which is a process that our body generates before various situations that have to do with the food that is ingested, solar radiation, pollution, and even excessive exercise. Oxidative stress plays a very important role in the development of diabetes complications due to excessive oxidative activity [23]. Just as there is stress that affects people with diabetes, there are also techniques that can be used to manage stress and maintain metabolic control, such as paying attention to one's breathing, social support, and progressive relaxation, among others [24].

2.2. Fuzzy Logic and Applications

Nowadays, systems that make use of fuzzy logic can be visualized in a more common way and can be seen in everyday life, such as in household appliances, industrial processes, medical measurement instruments, etc. Since fuzzy logic was proposed, we have witnessed that its applications have increased to this day. It was originally proposed in 1965 by Professor Zadeh, who presented the theory of fuzzy sets to deal with imprecision and uncertainty. A fuzzy set is defined as a class of objects that have degrees of membership, and these present the degree of truth with which a certain element belongs to a fuzzy set [25–27]. Different examples are presented to explain the fuzzy set concept, and many times, the word “fuzzy” itself is misunderstood in its meaning. A simple example is explained using the height of people: when it is said that a person is tall, for some, being tall refers to measuring more than 1.90 m, but for others, being tall refers to measuring more than 1.70 m. Each person has their perspective, so ranges can be managed; that is, membership functions such that, for example, height can be represented with more datasets (see Figure 1).

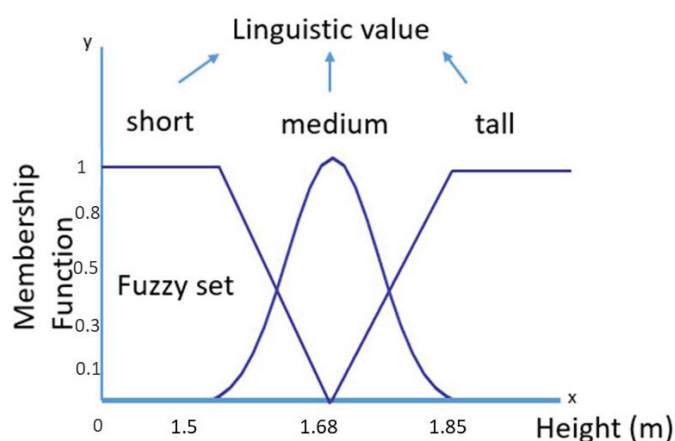


Figure 1. Fuzzy set representation.

In Figure 1, it can be seen that there are three linguistic values, namely short, medium, and tall, that can be used to classify people based on their height; for example, a person who is 1.68 m tall could be considered to be of average height, and a person who is 1.70 m could be in the range of consideration as a tall person but could also be considered as medium. This is just one example of how a fuzzy system would be considered, but it is also important to mention that one of the first commercial applications of fuzzy logic was in the control of cement kilns and later in navigation systems for automobiles [28]. One of the areas where fuzzy logic is implemented in a very competitive way is in the control area, making different types of controllers to improve the performance of applications and obtain a better result in the implementation.

Type-1 fuzzy systems are the most used, but other types of systems have also been found more recently, such as generalized systems and type-2 fuzzy systems, because the values used in the development of type-1 membership are considered to be accurate [29]. However, if the level of information is not adequate to establish the membership functions, then it would not work to establish the membership functions with this precision, but if it is regarding fuzzy systems that can be further generalized due to their interval values, these intervals could become fuzzy [30]. Therefore, each interval of the membership function would become an ordinary fuzzy set, and this is what is called an interval type-2 fuzzy set [31], which can be visualized in Figure 2.

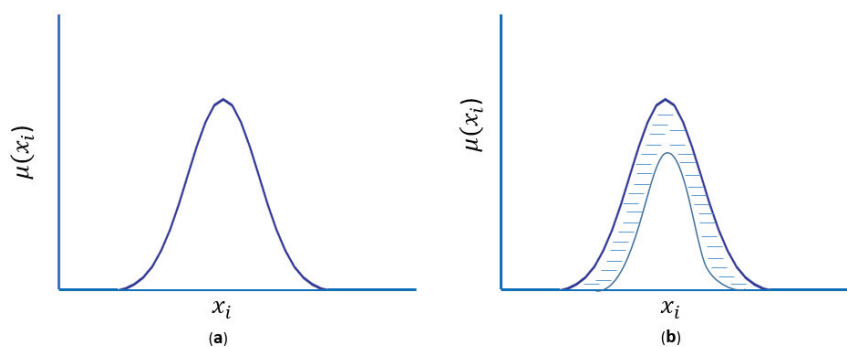


Figure 2. The membership function representation: (a) Type-1 fuzzy membership function; (b) type-2 fuzzy membership function.

Mathematical definition of membership function type-1 is obtained by in Equation (1):

$$A = \{(x, \mu_A(x)) \mid x \in X\} \quad (1)$$

Let us consider fuzzy set A , $A = \{(x, \mu_A(x)) \mid x \in X\}$, where $\mu_A(x)$ is called the membership function for the fuzzy set A . X is referred to as the universe. The membership function associates each element $x \in X$ with a value in the interval $[0, 1]$.

Mathematical definition of membership function type-2:

A type-2 fuzzy set to be named can be represented as follows [32] in Equation (2):

$$\tilde{A} = \{((x, u), \mu_{\tilde{A}}(x, u)) \mid \forall x \in X, \forall u \in J_x \subseteq [0, 1]\} \quad (2)$$

A type-2 fuzzy set (also called a generalized type-2 fuzzy set) denoted $\tilde{A}$, also called membership function of $\tilde{A}$, is about the product cartesian $X \times [0, 1]$ in $[0, 1]$, where X is the universe for the primary variable of $\tilde{A}$, x . the membership function of $\tilde{A}$ is represented as $\mu_{\tilde{A}}(x, u)$ or $\mu_{\tilde{A}}$, and it is called membership function type-2 (Equation (2)).

Fuzzy systems have had applications in several areas, and one of them is medicine; for example, in [33], the authors mention that there are different patients, and each one has different behavior patterns and coordination, especially in patients with diabetic cardiomyopathy. Since they have Parkinson's, they used a fuzzy system for the prediction of falls and the estimation of diabetic cardiomyopathy disorders, and as inputs, they took the height, gyroscope, age and weight of the patients. The outputs were the RMS error, estimation, and identification. The test was performed with 20 patients, and good results were obtained. The solution was fast and in real-time conditions. The fuzzy systems used were type-1, and in this example, we can also see fuzzy systems that are applied to patients to control their hypertension during anesthesia [34] as well as to diagnose certain types of diseases through these fuzzy systems and detection of cardiac arrhythmias [35,36].

In [37], studies developed from 2007 to 2018 are presented, including the current state of the use of fuzzy systems in the field of medicine for the diagnosis of diseases, and it is demonstrated that fuzzy logic has been used with satisfactory results. In the area of medicine, different methodologies are used to diagnose the disease by analyzing the history, symptoms, and clinical data of people. The study presents the benefits of using fuzzy logic for this area, which is an important detail to mention. These systems could be available to people, and some important issues could be identified before going directly to the doctor or even helping through medicine to diagnose diseases. Some areas within medicine where fuzzy systems have been applied are in the heart, breast cancer, cholera, brain tumor, asthma, liver, viral diseases, etc. [38–40]. Many applications begin with developing studies with type-1 fuzzy systems but when, for example, the images or databases have a high level of noise or uncertainty, the outcome can be unsure. In these cases, type-2 fuzzy systems could work better and obtain better performance and model better the uncertainty.

Some studies in medicine also use fuzzy type-2 systems. For example, in [41], the authors present a proposal to design a robust and optimal controller for the supply of medicine for the automatic control of blood pressure, as all human systems also have uncertainty problems, such as variations in certain parameters, disturbances, or external noises. That is why for each patient with different characteristics, an adequate supply of medicine is needed for good control. On many occasions, PID controllers have been used, even in controllers for health at a commercial level, but in some cases, with the management of uncertainty, a more robust control would be needed that adapts to the characteristics of each individual since all the physiological variables that a person may have are of an uncertain nature and can negatively affect their health, which is why in these cases tests are carried out with type-1 fuzzy systems and then type-2 [42] to better model the uncertainty. In the area of diabetes, research papers have also emerged, and everything possible is done to contribute to scientific knowledge and provide work ideas so that in the future, these ideas are taken to improve applications. In [43], the authors present a fuzzy system to classify the glycemic index for patients with type 2 diabetes into four types: great decrease, decreases followed by stabilizing, stabilizing, and then increasing. In this research, some authors such as in [44,45] also performed a fuzzy system for glucose control in diabetic patients. Many intelligent techniques are applied in several health areas to improve the quality of people's lives [46–50].

3. Problem Description

As previously mentioned, type 1 diabetes is a disease that can have serious health complications if it is not treated properly. Research has been carried out in relation to this disease, and controllers have even been designed for the delivery of the necessary insulin. However, there are many different types of extra problems unique to each individual, such as different situations involving stress or some other external problem that can influence their glucose levels. In this situation, a common controller will not be able to completely regulate what we call uncertainty or noise, which can represent the different levels of stress experienced by people or the external factors that influence each one of them, affecting their glucose levels in real time. The human, as mentioned above, can experience a great deal of uncertainty, which in this case will be called noise, which will be added to the system as an extra signal. The real-time controller will model the noise, and in the event of glucose disturbances, the insulin pump can supply the optimal dose needed to maintain control of the diabetes with the aim of improving people's quality of life.

The human body generates signals before any behavior, and in this case, the noise is represented by the signal established by the following equation:

$$y = lb + (ub - lb) \cdot \text{rand}(spf, 1) \quad (3)$$

The working model can be seen in Figure 3.

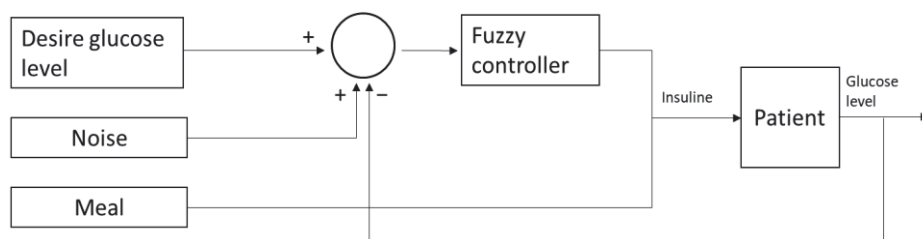


Figure 3. Diagram for the insulin control.

The development model contemplates a module named noise, which is represented by a module that generates a signal that can be small or very high so that the system models the adequate supply of insulin in different levels of uncertainty that the patient may have.

Three types of error were considered: root mean square error (*RMSE*), mean square error (*MSE*), and mean absolute error (*MAE*) (see Equations (4)–(6)).

$$RMSE = \sqrt{\frac{1}{N} \sum_{i=1}^N (real\ value_i - estimated\ value_i)^2} \quad (4)$$

$$MSE = \frac{1}{N} \sum_{i=1}^N (real\ value_i - estimated\ value_i)^2 \quad (5)$$

$$MAE = \frac{1}{N} \sum_{i=1}^N |real\ value_i - estimated\ value_i| \quad (6)$$

where *N* is the number of errors, and before using the fuzzy systems, the performance was simulated using PID control; then, the type-1 and type-2 fuzzy systems were designed to compared results and to achieve the control of the insulin in each person. The fuzzy system has one input and one output.

4. Simulation Results

As the first execution, PID was performed, and then, different types of fuzzy systems were design to test the model using three or five membership functions. By changing the rules, 15 different fuzzy inference systems (fis) were developed, and results for the PID controller and the 15 fuzzy systems are presented in Table 1.

Table 1. Results for the PID and type-1 fuzzy control for optimal supply of insulin.

Controller	RMSE Error	MAE Error	MSE Error
PID	21.0947	17.8862	444.9858
Type-1			
Fuzzy controller	RMSE error	MAE error	MSE error
fis 1	28.810	27.560	830.003
fis 2	20.347	15.710	413.990
fis 3	26.973	21.559	727.514
fis 4	19.149	14.562	366.694
fis 5	18.271	13.746	333.821
fis 6	17.640	13.150	311.176
fis 7	19.318	14.750	373.199
fis 8	19.745	15.144	389.882
fis 9	16.783	15.180	281.679
fis 10	11.883	8.950	141.201
fis 11	11.874	8.879	140.979
fis 12	11.912	8.960	141.892
fis 13	11.871	8.908	140.911
fis 14	11.955	8.975	142.929
fis 15	15.655	11.375	245.088
Average	17.479	13.827	332.064

When the type-1 fuzzy system was developed with different membership functions and rules, the error decreased until 11.871 with the fis 13. The parameters of the best fuzzy system used to achieve the control of the insulin pump injection are presented in Figure 4, and rules are illustrated in Table 2. In this case, the linguistic values for the variables used in the set rules are VL, very low; L, low; M, medium; H, high; and VH, very high. The number of rules is the combination of the input and output membership functions. It is important to mention to the reader that the fis number of 13 was the one that obtained the best result for the 15 tests carried out.

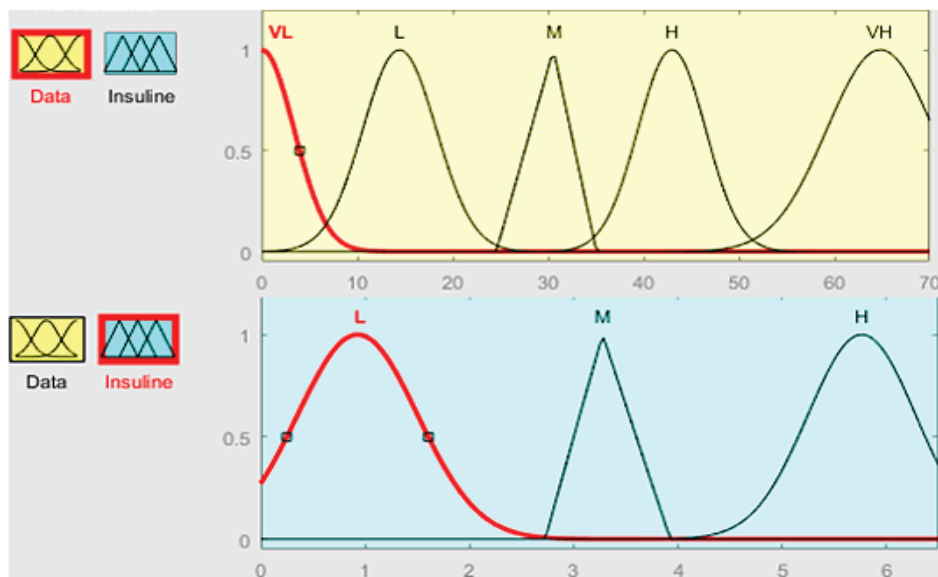


Figure 4. Best type-1 fuzzy system to control the insulin pump injection.

Table 2. Rules used in the type-1 fuzzy system.

Input	Output
VL	L
VL	H
VL	M
M	L
M	H
M	M
VH	L
VH	H
VH	M
L	L
L	H
L	M
H	L
H	H
H	M

The parameters of the membership functions are presented in Table 3.

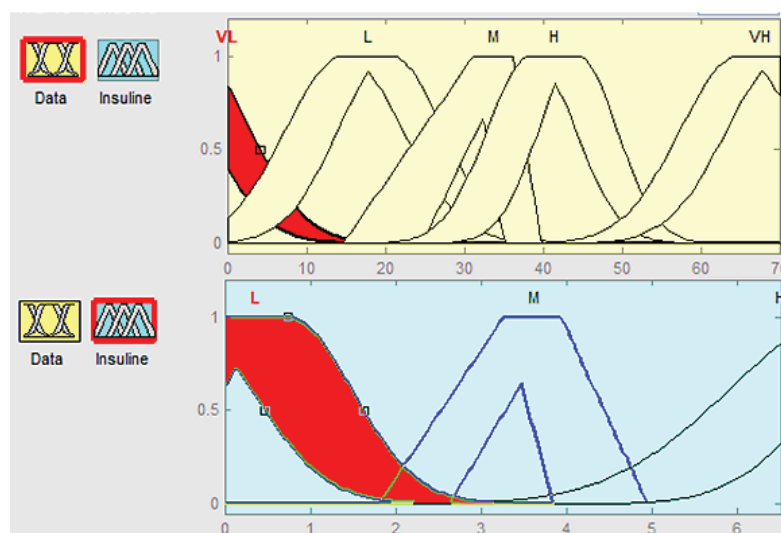
Table 3. Parameters of the type-1 membership functions.

	Linguistic Variable	Membership Function Type	Values
Input	VL	Gaussian	3.382 1.01×10^{-15}
	L	Gaussian	3.85 14.45
	M	Triangular	24.56 30.58 35.1
	H	Gaussian	3.542 42.97
	VH	Gaussian	5.586 64.81
Output	L	Gaussian	0.5763 0.9286
	M	Triangular	2.726 3.276 3.926
	H	Gaussian	0.5174 5.761

After obtaining results with the type-1 fuzzy systems, the type-2 fuzzy system tests were performed, and results are shown in Table 4. The best type-2 fuzzy system is presented in Figure 5, and parameters of each membership functions are presented in Table 5.

Table 4. Results for the type-2 fuzzy control for optimal supply of insulin.

Type-2 Fuzzy Controller	RMSE Error	MAE Error	MSE Error
fis1	10.948	12.592	119.853
fis2	6.638	8.282	44.064
fis3	7.655	9.265	58.605
fis4	6.319	8.601	39.935
fis5	7.143	9.425	51.028
fis6	6.834	9.116	46.709
fis7	7.391	9.673	54.633
fis8	5.621	11.235	31.592
fis9	8.293	14.543	68.767
fis10	5.060	16.941	25.606
fis11	2.318	10.199	5.374
fis12	4.854	6.498	23.564
fis13	5.854	7.498	34.272
fis14	6.081	7.725	36.977
fis15	4.274	5.918	18.267
Average	6.352	9.834	43.950

**Figure 5.** Best type-2 fuzzy system to control the insulin pump injection.**Table 5.** Parameters of the type-2 membership functions.

	Linguistic Variable	Membership Function Type	Values
Input	VL	Gaussian	7.43 –10.1 7.114 –4.21
	L	Gaussian	7.43 14.8 6.649 20.7
	M	Triangular	14.6 31.2 35.1 24.7 36.2 39.722
	H	Gaussian	5.6 38.529 4.99 44.429
	VH	Gaussian	7.43 64.887 7.43 70.787
Output	L	Gaussian	0.862 –0.542 0.759 0.731
	M	Triangular	1.791 3.258 3.841 2.621 3.931 4.941
	H	Gaussian	1.408 7.29 1.02 8.03

After obtaining results using type-1 and type-2 fuzzy systems with three different types of error, a statistical test was performed as presented in Tables 6–8.

Table 6. Results for Student's *t*-test using RMSE error in type-1 and type-2 fuzzy controllers.

Controller	N	Mean	Deviation Std.	Mean of Std. Error
Type-1	15	17.48	5.33	1.4
Type-2	15	6.35	1.96	0.51
T value = 7.58; <i>p</i> -value = 0.000; DF = 28				

Table 7. Results for Student's *t*-test using MAE error in type-1 and type-2 fuzzy controllers.

Controller	N	Mean	Deviation Std.	Mean of Std. Error
Type-1	15	13.83	5.23	1.4
Type-2	15	9.83	2.97	0.77
T value = 2.57; <i>p</i> -value = 0.016; DF = 28				

Table 8. Results for Student's *t*-test using MSE error in type-1 and type-2 fuzzy controllers.

Controller	N	Mean	Deviation std.	Mean of Std. Error
Type-1	15	332	209	54
Type-2	15	43.9	26.7	6.9
T value = 5.31; <i>p</i> -value = 0.000; DF = 28				

As can be seen in Tables 4–6, there is statistical evidence to say that there is a significant difference in the results presented although we used different types to calculate the error with regard to type-2 (more than 95% confidence). Basically, type-2 fuzzy systems generated a significant increase in the controller, enabling better control for the insulin dose in response to disturbances.

The statistical test was a critical point to test the insulin control supply in real time using fuzzy systems when the type-1 fuzzy system the minimum error was 11.87, as mentioned above. Taking into account that people's health is very important, system with greater robustness, such as the fuzzy type-2 system, was developed, which in many cases can give better results because it has better response in uncertainty. When we performed the tests with the type-2 fuzzy system, the behavior of the control improved, and we can verify it with the result of the statistical tests that were carried out previously.

5. Conclusions

Health is one of the most valuable things that human beings have because if they are not healthy, in addition to affecting their personal health, it can have a negative impact on the people around them. The importance of having an optimal control of disease is paramount; in this case, a person with type 1 diabetes can stop worrying about their optimal dose supplied by the insulin pump in case they have a stressful situation or some external or even internal factor that could affect their glucose levels, which would improve the emotional condition of the patient and family members. Fuzzy systems, thanks to their versatility, are adaptable to several characteristics and situations that the patient may present. In this case, it was verified that type-2 fuzzy systems for insulin supply pump control work better than type-1 fuzzy systems and the PID since the patient is exposed to situations where noise levels could occur at low levels or at high levels of uncertainty, and the fuzzy model type-2 helps in these cases. In the last section of the paper, a statistical Student's *t*-test was performed to prove the efficiency of the type-1 and type-2

fuzzy systems, and results showed that a type-2 fuzzy system can model the uncertainty in a better way to obtain the best control of insulin using pump injection in real time. The statistical comparison was a very important part of the results, as it demonstrates a significant difference in real-time insulin control. These results show with more than 95% confidence that using type-2 fuzzy systems to control insulin supply works better, and people can have an optimal insulin supply in real time, with no need to manually check. The statistical test was required to be able to numerically visualize the operation of type-1 and type-2 fuzzy systems and verify which is the best for this type of case study.

Author Contributions: Conceptualization, L.C. and C.C.; methodology, L.C.; software, C.C.; validation, L.C., C.C. and O.C.; formal analysis, O.C.; investigation, C.C. and L.C.; writing—review and editing, L.C., O.C. and C.C. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Data Availability Statement: Not applicable.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Association, A.D. Classification and Diagnosis of Diabetes: Standards of Medical Care in Diabetes—2020. *Diabetes Care* **2019**, *43* (Suppl. S1), S14–S31. [CrossRef] [PubMed]
2. Ali, W.G.; Nagib, G. Embedded Control Design for Insulin Pump. In *Advanced Manufacturing Systems, ICMSE 2011*; Advanced Materials Research; Trans Tech Publications Ltd.: Stafa-Zurich, Switzerland, 2011; Volume 201, pp. 2399–2404. [CrossRef]
3. Karges, B.; Schwandt, A.; Heidtmann, B.; Kordonouri, O.; Binder, E.; Schierloh, U.; Boettcher, C.; Kapellen, T.; Rosenbauer, J.; Holl, R.W. Association of Insulin Pump Therapy vs. Insulin Injection Therapy With Severe Hypoglycemia, Ketoacidosis, and Glycemic Control Among Children, Adolescents, and Young Adults With Type 1 Diabetes. *JAMA* **2017**, *318*, 1358–1366. [CrossRef]
4. Bergenstal, R.M.; Nimri, R.; Beck, R.W.; Criego, A.; Laffel, L.; Schatz, D.; Battelino, T.; Danne, T.; Weinzimer, S.A.; Sibayan, J.; et al. A Comparison of Two Hybrid Closed-Loop Systems in Adolescents and Young Adults with Type 1 Diabetes (FLAIR): A Multicentre, Randomised, Crossover Trial. *Lancet* **2021**, *397*, 208–219. [CrossRef]
5. Álvarez Casaño, M.; del Mar Alonso Montejó, M.; Leiva Gea, I.; Jiménez Hinojosa, J.M.; Santos Mata, M.Á.; Macías, F.; del Mar Romero Pérez, M.; de Toro, M.; Martínez, G.; Munguira, P.; et al. Estudio de Calidad de Vida y Adherencia Al Tratamiento En Pacientes de 2 a 16 Años Con Diabetes Mellitus Tipo 1 En Andalucía. *An. Pediatría* **2021**, *94*, 75–81. [CrossRef] [PubMed]
6. DeMarsilis, A.; Mantzoros, C.S. The Continuum of Insulin Development Viewed in the Context of a Collaborative Process toward Leveraging Science to Save Lives: Following the Trail of Publications and Patents One Century after Insulin's First Use in Humans. *Metabolism* **2022**, *135*, 155251. [CrossRef] [PubMed]
7. Díez Gutiérrez, B. Curso Básico Sobre Diabetes. Tema 1. Clasificación, Diagnóstico y Complicaciones. *Farm. Prof.* **2016**, *30*, 36–43. [CrossRef]
8. Saldarriaga-Giraldo, C.; Navas, V.; Morales, C. De La Diabetes a La Insuficiencia Cardíaca ¿Existe La Miocardiopatía Diabética? *Rev. Colomb. Cardiol.* **2020**, *27*, 12–16. [CrossRef]
9. Tittel, S.R.; Rosenbauer, J.; Kamrath, C.; Ziegler, J.; Reschke, F.; Hammersen, J.; Mönkemöller, K.; Pappa, A.; Kapellen, T.; Holl, R.W. Did the COVID-19 Lockdown Affect the Incidence of Pediatric Type 1 Diabetes in Germany? *Diabetes Care* **2020**, *43*, e172–e173. [CrossRef]
10. Alonso, N.; Batule, S. COVID-19 y Diabetes Mellitus. Importancia Del Control Glucémico. *Clínica E Investig. En Arterioscler.* **2021**, *33*, 148–150. [CrossRef]
11. Hernández Herrero, M.; Terradas Mercader, P.; Latorre Martínez, E.; Feliu Rovira, A.; Rodríguez Zaragoza, N.; Parada Ricart, E. Nuevos Diagnósticos de Diabetes Mellitus Tipo 1 En Niños Durante La Pandemia COVID-19. Estudio Multicéntrico Regional En España. *Endocrinol. Diabetes Y Nutr.* **2022**, *69*, 709–714. [CrossRef]
12. Dovc, K.; Battelino, T. Evolution of Diabetes Technology. *Endocrinol. Metab. Clin. North Am.* **2020**, *49*, 1–18. [CrossRef]
13. Barrio Castellanos, R. Avances En El Tratamiento de La Diabetes Tipo 1 Pediátrica. *An. Pediatría* **2021**, *94*, 65–67. [CrossRef] [PubMed]
14. Huard, B.; Kirkham, G. Mathematical Modelling of Glucose Dynamics. *Curr. Opin. Endocr. Metab. Res.* **2022**, *25*, 100379. [CrossRef]
15. Poddar, A.; Rangwani, N.; Palekar, S.; Kalambe, J. Glucose Monitoring System Using Machine Learning. *Mater. Today Proc.* **2022**, *73*, 100–107. [CrossRef]
16. Haynes, E.; Ley, M.; Talbot, P.; Dunbar, M.; Cummings, E. Insulin Pump Therapy Improves Quality of Life of Young Patients With Type 1 Diabetes Enrolled in a Government-Funded Insulin Pump Program: A Qualitative Study. *Can. J. Diabetes* **2021**, *45*, 395–402. [CrossRef] [PubMed]

17. Chandrasekhar, A.; Padhi, R. Optimization Techniques for Online MPC in Android Smartphones for Artificial Pancreas: A Comparison Study. *IFAC-PapersOnLine* **2022**, *55*, 561–566. [CrossRef]
18. Farman, M.; Saleem, M.U.; Ahmad, A.; Imtiaz, S.; Tabassum, M.F.; Akram, S.; Ahmad, M.O. A Control of Glucose Level in Insulin Therapies for the Development of Artificial Pancreas by Atangana Baleanu Derivative. *Alexandria Eng. J.* **2020**, *59*, 2639–2648. [CrossRef]
19. Aguilar, M.E.B. El Estrés y Su Influencia En La Calidad de Vida. *Multimed* **2017**, *21*, 971–982.
20. Alberto, R.S.C. *Procesos Complejos Del Estrés: Dinámica No-Lineal*; Universidad Nacional de Colombia: Bogotá, Colombia, 2012.
21. Beléndez Vázquez, M.; Lorente Armendáriz, I.; Maderuelo Labrador, M. Estrés Emocional y Calidad de Vida En Personas Con Diabetes y Sus Familiares. *Gac. Sanit.* **2015**, *29*, 300–303. [CrossRef]
22. Juárez Jiménez, M.V. Influencia Del Estrés En La Diabetes Mellitus. *NPunto* **2020**, *3*, 91–124.
23. Storino, M.A.; Contreras, M.A.; Rojano, J.; Serrano, R.; Nouel, A. Complicaciones de La Diabetes y Su Asociación Con El Estrés Oxidativo: Un Viaje Hacia El Daño Endotelial. *Rev. Colomb. Cardiol.* **2014**, *21*, 392–398. [CrossRef]
24. Montes Delgado, R.; Oropeza Tena, R.; Pedroza Cabrera, F.J.; Verdugo Lucero, J.C.; Enríquez Bielma, J.F. Manejo Del Estrés Para El Control Metabólico de Personas Con Diabetes Mellitus Tipo 2. *En-Claves Del Pensam.* **2013**, *7*, 67–87.
25. Simić, D.; Kovačević, I.; Svirčević, V.; Simić, S. 50 Years of Fuzzy Set Theory and Models for Supplier Assessment and Selection: A Literature Review. *J. Appl. Log.* **2017**, *24*, 85–96. [CrossRef]
26. Lotfi, A.; Zadeh, R.A.A. *Fuzzy Logic Theory and Applications: Part I and Part I*; World Scientific Publishing: Singapore, 2018. [CrossRef]
27. Höhle, U. On the Mathematical Foundations of Fuzzy Set Theory. *Fuzzy Sets Syst.* **2022**, *444*, 1–9. [CrossRef]
28. Voskoglou, M. *Fuzzy Sets, Fuzzy Logic and Their Applications*; MDPI-Multidisciplinary Digital Publishing Institute: Basel, Switzerland, 2020. [CrossRef]
29. Mendel, J.M. *Uncertain Rule-Based Fuzzy Systems Introduction and New Directions*, 2nd ed.; Springer: Cham, Switzerland, 2017. [CrossRef]
30. Timothy, J.R. *Fuzzy Logic with Engineering Applications*, 3rd ed.; Wiley: New York, NY, USA, 2010.
31. Castillo, O.; Aguilar, L.T. *Background on Type-1 and Type-2 Fuzzy Logic BT—Type-2 Fuzzy Logic in Control of Nonsmooth Systems: Theoretical Concepts and Applications*; Castillo, O., Aguilar, L.T., Eds.; Springer International Publishing: Cham, Switzerland, 2019; pp. 5–19. [CrossRef]
32. Najariyan, M.; Mazandarani, M.; John, R. Type-2 Fuzzy Linear Systems. *Granul. Comput.* **2017**, *2*, 175–186. [CrossRef]
33. Sharma, N.; Sampath Dakshina Murthy, A.; Karthikeyan, T.; Usha Kumari, C.; Omkar Lakshmi Jagan, B. Gait Diagnosis Using Fuzzy Logic with Wearable Tech for Prolonged Disorders of Diabetic Cardiomyopathy. *Mater. Today Proc.* **2020**, online. [CrossRef]
34. Oshita, S.; Nakakimura, K.; Sakabe, T. Hypertension Control during Anesthesia. Fuzzy Logic Regulation of Nicardipine Infusion. *IEEE Eng. Med. Biol. Mag.* **1994**, *13*, 667–670. [CrossRef]
35. Kora, P.; Meenakshi, K.; Swaraja, K.; Rajani, A.; Kafiul Islam, M. Detection of Cardiac Arrhythmia Using Fuzzy Logic. *Inform. Med. Unlocked* **2019**, *17*, 100257. [CrossRef]
36. Iancu, I. Heart Disease Diagnosis Based on Mediative Fuzzy Logic. *Artif. Intell. Med.* **2018**, *89*, 51–60. [CrossRef]
37. Thukral, S.; Rana, V. Versatility of Fuzzy Logic in Chronic Diseases: A Review. *Med. Hypotheses* **2019**, *122*, 150–156. [CrossRef]
38. Kumar, A.; Raj, R. Design of a Fractional Order Two Layer Fuzzy Logic Controller for Drug Delivery to Regulate Blood Pressure. *Biomed. Signal Process. Control* **2022**, *78*, 104024. [CrossRef]
39. Thakkar, H.; Shah, V.; Yagnik, H.; Shah, M. Comparative Anatomization of Data Mining and Fuzzy Logic Techniques Used in Diabetes Prognosis. *Clin. eHealth* **2021**, *4*, 12–23. [CrossRef]
40. Arji, G.; Ahmadi, H.; Nilashi, M.; Rashid, T.A.; Ahmed, O.H.; Aljojo, N.; Zainol, A. Fuzzy Logic Approach for Infectious Disease Diagnosis: A Methodical Evaluation, Literature and Classification. *Biocybern. Biomed. Eng.* **2019**, *39*, 937–955. [CrossRef]
41. Sharma, R.; Deepak, K.K.; Gaur, P.; Joshi, D. An Optimal Interval Type-2 Fuzzy Logic Control Based Closed-Loop Drug Administration to Regulate the Mean Arterial Blood Pressure. *Comput. Methods Programs Biomed.* **2020**, *185*, 105167. [CrossRef] [PubMed]
42. Doctor, F.; Syue, C.-H.; Liu, Y.-X.; Shieh, J.-S.; Iqbal, R. Type-2 Fuzzy Sets Applied to Multivariable Self-Organizing Fuzzy Logic Controllers for Regulating Anesthesia. *Appl. Soft Comput.* **2016**, *38*, 872–889. [CrossRef]
43. Bressan, G.M.; Azevedo, B.F.; de Souza, R.M. A Fuzzy Approach for Diabetes Mellitus Type 2 Classification. *Brazilian Arch. Biol. Technol.* **2020**, *63*. [CrossRef]
44. Zhu, K.Y.; Liu, W.D.; Xiao, Y. *Application of Fuzzy Logic Control for Regulation of Glucose Level of Diabetic Patient BT-Machine Learning in Healthcare Informatics*; Dua, S., Acharya, U.R., Dua, P., Eds.; Springer: Berlin/Heidelberg, Germany, 2014; pp. 47–64. [CrossRef]
45. Asadi, S.; Nekoukar, V. Adaptive Fuzzy Integral Sliding Mode Control of Blood Glucose Level in Patients with Type 1 Diabetes: In Silico Studies. *Math. Biosci.* **2018**, *305*, 122–132. [CrossRef] [PubMed]
46. Sharma, A.; Nilam; Singh, H.P. Computer-Controlled Diabetes Disease Diagnosis Technique Based on Fuzzy Inference Structure for Insulin-Dependent Patients. *Appl. Intell.* **2023**, *53*, 1945–1958. [CrossRef]
47. Sohrabi, M.; Zandieh, M.; Shokouhifar, M. Sustainable Inventory Management in Blood Banks Considering Health Equity Using a Combined Metaheuristic-Based Robust Fuzzy Stochastic Programming. *Socioecon. Plann. Sci.* **2022**, 101462. [CrossRef]
48. Ghorbani, A.; Davoodi, F.; Zamanifar, K. Using Type-2 Fuzzy Ontology to Improve Semantic Interoperability for Healthcare and Diagnosis of Depression. *Artif. Intell. Med.* **2023**, *135*, 102452. [CrossRef]

49. Xiao, G.; Xiao, Y.; Ni, A.; Zhang, C.; Zong, F. Exploring influence mechanism of bikesharing on the use of public transportation—A case of Shanghai. *Transp. Lett.* **2022**, 1–9. [CrossRef]
50. Chen, X.; Wu, S.; Shi, C.; Huang, Y.; Yang, Y.; Ke, R.; Zhao, J. Sensing data supported traffic flow prediction via denoising schemes and ANN: A comparison. *IEEE Sens. J.* **2020**, 20, 14317–14328. [CrossRef]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

Article

A Fuzzy-Based Fast Feature Selection Using Divide and Conquer Technique in Huge Dimension Dataset

Arihant Tanwar¹, Wajdi Alghamdi^{2,*}, Mohammad D. Alahmadi³, Harpreet Singh¹
and Prashant Singh Rana¹

¹ Department of Computer Science and Engineering, Thapar Institute of Engineering and Technology, Patiala 147004, Punjab, India; arihant.tanwar21@gmail.com (A.T.); harpreet.s@thapar.edu (H.S.); prashant.singh@thapar.edu (P.S.R.)

² Department of Information Technology, Faculty of Computing and Information Technology, King Abdulaziz University, Jeddah 21589, Saudi Arabia

³ Department of Software Engineering, College of Computer Science and Engineering, University of Jeddah, Jeddah 21959, Saudi Arabia; mdalahmadi@uj.edu.sa

* Correspondence: wmalghamdi@kau.edu.sa

Abstract: Feature selection is commonly employed for identifying the top n features that significantly contribute to the desired prediction, for example, to find the top 50 or 100 genes responsible for lung or kidney cancer out of 50,000 genes. Thus, it is a huge time- and resource-consuming practice. In this work, we propose a divide-and-conquer technique with fuzzy backward feature elimination (FBFE) that helps to find the important features quickly and accurately. To show the robustness of the proposed method, it is applied to eight different datasets taken from the NCBI database. We compare the proposed method with seven state-of-the-art feature selection methods and find that the proposed method can obtain fast and better classification accuracy. The proposed method will work for qualitative, quantitative, continuous, and discrete datasets. A web service is developed for researchers and academicians to select top n features.

Keywords: feature selection; divide-and-conquer technique; huge dimension dataset; genomic dataset; fuzzy technique; fuzzy backward feature elimination

MSC: 15-04

1. Introduction

A feature is an individual measurable property of the process being observed. A machine learning algorithm predicts the value of the desired target variable using these features [1]. We are now in the era of big data, where huge amounts of high-dimensional data have become ubiquitous in various domains, such as social media, healthcare, bioinformatics, and online education. The rapid growth of data presents challenges for feature selection. The “curse of dimensionality” (CoD), a wealth of riches, presents itself in various forms [2]. Feature Selection (variable elimination) helps understand the data, reduce computation requirements, reduce the effect of dimensionality’s curse, and improve the predictor performance.

Let, ρ = number of observed variables.

Initially, as the dimensionality ρ rises, the space that the samples could occupy expands rapidly. Figure 1 shows the model performance with respect to number for features. If we consider the distance between the points as a measure of similarity, then we interpret the greater distance as a greater dissimilarity. As ρ increases, the pairwise distance between two points decreases, the correlation among vectors increases, and the likelihood of a specific region of the space being empty and sparse, with no data, increases.

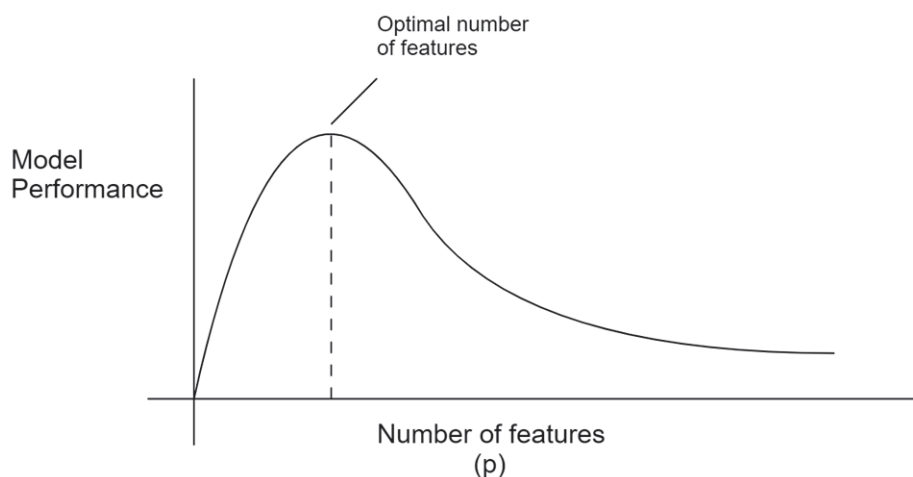


Figure 1. Effect of number of features on performance of model.

In other words, as the number of dimensions grows, the amount of data required for satisfactory results from any machine learning algorithm expands rapidly. The cause of this is that with more dimensions, the model needs a greater number of data points to represent all the possible combinations of features in order to be considered valid.

Hughes (1968) found in his study that the effectiveness of a classifier in predicting outcomes improves as the number of dimensions increases up to a point [3]. Beyond that point, the performance decreases. This phenomenon of diminishing returns in prediction accuracy as the number of dimensions grows is commonly referred to as the “curse of dimensionality” or the “Hughes phenomenon”.

This study shows the requirement of feature selection, especially for a huge dimension dataset. Existing feature selection methods suffice for small datasets but fail in huge dimension datasets due to high computational requirements. A feature selection method removes features by measuring the relevance of each feature with the target class or by measuring the correlation between features. If two features are highly correlated, then one of them is removed. Feature selection methods are different from dimension reduction methods, such as PCA [1,4]. This is because good features can be independent from the rest of the data [5].

The outcome of a feature selection attempt from a broad dataset depends on several factors, such as the underlying probability distributions (some issues could be simple to solve), the number of instances (sample size), the number of features (dimensionality), the selected method for feature selection (its ability to find optimal feature subsets, its resistance to overfitting, and the accuracy of evaluating the desired criterion), and the classifier recommended to the user, as noted in [6].

Kohavil et al. propose the wrappers method for feature subset selection, which is divided into two major categories: filter and wrapper methods [7]. Filter methods score features without testing them in prediction algorithms, while in wrapper methods, the feature selection criteria are as per the predictor’s performance. The embedded method [8] is another technique that includes a variable method as a part of the training process.

2. Feature Selection Methods

2.1. Filter Methods

Filter methods choose features based on a measure of performance, independent of the data modeling algorithm used. The algorithm uses those features that are selected from filter methods. Filter methods can evaluate individual features or entire feature subsets, ranking them based on performance. These methods can be applied to various problems in classification, clustering, and regression, according to [9]. These criteria are mutual information, correlation, f-score, and chi-square.

2.1.1. Mutual Information

This method evaluates the dependency between two variables to score each of them. To understand mutual information [10], we must start with an entropy of variables.

Let X = random variable with discrete values. Then, the entropy of X , $H(X)$ can be defined as

$$H(X) = - \sum p(x) \log p(x), \quad x \in X \quad (1)$$

where $p(x)$ is the probability density function of x .

Conditional entropy refers to the uncertainty reduction of a variable when another is known. Assuming that variable Y is given, the conditional entropy $H(X|Y)$ of X with respect to Y is

$$H(X|Y) = - \sum \sum p(x, y) \log p(x|y), \quad x \in X, y \in Y \quad (2)$$

Equation (2) indicates that observing the variable X reduces the uncertainty in Y . The decrease in uncertainty is expressed as

$$I(Y, X) = H(Y) - H(Y|X) \quad (3)$$

This gives the mutual information between X and Y . Mutual information between X and Y will be zero if they are independent; greater than zero, they are dependent.

2.1.2. Correlation

Correlation is a statistical measure expressed as a number that characterizes the magnitude and direction of the relationship between two or more variables. The most commonly used measure of dependence between two variables is the Pearson correlation coefficient [11], which can be defined as

$$\rho_{(X,Y)} = \text{corr}(X, Y) = \frac{\text{cov}(X, Y)}{\sigma_X \sigma_Y} = \frac{E[(X - \mu_X)(Y - \mu_Y)]}{\sigma_X \sigma_Y}, \quad \text{if } \sigma_X \sigma_Y > 0 \quad (4)$$

The above equation calculates the correlation between X and Y . The $\text{cov}(X, Y)$ calculates the covariance. The correlation ranking detects the correlation between features and the target variable.

2.1.3. Chi-Squared

A chi-squared test (symbolically represented as χ^2) is a data analysis based on observations of a random set of variables. The chi-squared method evaluates the association of two categorical variables. The importance of a feature increases if its chi score is high [12]. Chi-squared statistics can be defined as follows:

$$\chi^2 = \sum_{i=1}^m \sum_{j=1}^k \left(\frac{A_{ij} - E_{ij}}{E_{ij}} \right)^2 \quad (5)$$

where, m is the number of intervals, k is the number of classes, A_{ij} is the number of samples in the i th interval of the j th class, R_i is the number of samples in the i th interval, C_j is the number of samples in the j th class, N is the total number of samples, and is the expected frequency of A_{ij} ($E_{ij} = R_i * C_j / N$).

2.2. Wrapper Method

Wrapper methods employ a specific learning algorithm to assess the accuracy performance of a potential feature subset, leading to improved solutions [13]. In wrapper methods, a model is trained using a subset of features. Based on the performance of the model on these features, features are either added or removed from the subset. Some popular wrapper methods include forward selection, backward elimination, exhaustive feature selection, recursive feature elimination, and recursive feature elimination with cross-validation. Forward selection methods start with zero features and add up features

according to relevance. After this, the final set of selected features is returned. Elimination methods start with the set of all features and eliminate features in every iteration until the desired number of features is obtained.

However, as this method requires training of the model on every iteration, it has a high computational cost [14]. As the number of features grows, the space of feature subsets grows exponentially. This becomes critical when tens of thousands of features are considered, for example, in a genomics dataset. As a result, the wrapper approach is largely avoided.

2.3. Embedded Method

Embedded methods are a combination of filter and wrapper methods that are iterative in nature. They differ from other feature selection methods, as they involve the incorporation of feature selection as a part of the training process [8]. During each iteration of the model training process, embedded methods carefully extract features that contribute the most to the training. Unlike wrapper methods, the link between feature selection and classification algorithms is stronger in embedded methods, as they use classification algorithms that contain their own built-in ability to select features [15].

3. Related Work

Filter methods in feature selection can be divided into two categories: feature weighting and subset search algorithms. Feature weighting algorithms determine which features are most important by assigning scores to individual features and ranking them based on these scores. On the other hand, subset search algorithms evaluate the goodness of entire feature subsets and choose the best one according to a certain evaluation measure [16].

A study merged two feature selection/extraction algorithms, independent component analysis (ICA) and fuzzy backward feature elimination (FBFE), and applied them to five DNA microarray datasets [17].

In a study by Yasmin et al., a graph-based feature selection approach was presented for language identification using rough-set boundary regions [18]. A new system was proposed that leverages the rough set theory to enhance the accuracy of language identification by using roughness from the theory to construct a weighted graph.

Reimann et al. tackle real-world-sized vehicle routing problems through their research [19]. They evaluate both established benchmark instances and new, larger-scale vehicle routing problem instances. The authors show that their approach not only improves efficiency, but also increases the algorithm's effectiveness, leading to a highly effective tool for resolving real-world-sized vehicle routing problems.

Song et al. introduce a fast clustering-based feature subset selection algorithm that is designed for high-dimensional data [20]. The FAST algorithm consists of two phases. In the first phase, graph-based methods are utilized to classify features into clusters. In the next stage, the most significant features that have a strong association with the target classes are chosen from each cluster to form a subset of features.

Zhao et al. developed a recursive divide-and-conquer approach for sparse principal component analysis [21]. The approach divides the complex problem of sparse PCA into simpler sub-problems with known solutions and then solves each sub-problem recursively, resulting in a highly efficient algorithm for sparse PCA.

In the field of improving classification accuracy, multiple techniques have been proposed that aim to assign a shared discriminative feature set to the local behavior of data in different parts of the feature space [22,23]. One such method is localized feature selection (LFS), introduced by Armanfard et al. in which a subset of features is selected to fit a specific group of samples [23].

The class label of a new sample is assigned based on its similarity to the representative sample of each region in the feature space. The similarity is calculated for the sample's arbitrary query. Some feature selection methods rank features by using aggregated sample data, such as in the case of the approaches introduced by Tibshirani et al. and Chen et al. [24,25].

Some of the feature selection approaches that would be used for comparison with the proposed feature selection methods are discussed below:

3.1. Minimum Redundancy Maximal Relevance Criteria (mRMR)

In mRMR, the maximum dependency condition (mutual information) is transformed into an equivalent form for incremental feature selection at the first order. This is followed by the application of a two-stage feature selection algorithm which merges mRMR with feature selection techniques, such as the wrapper method, leading to a cost-effective feature selection process [26].

3.2. Least Angle Regression (LARS)

In “least angle regression (LARS)”, three key properties are established [27]. The LARS algorithm is a less aggressive version of traditional forward selection methods and can be modified in three ways:

1. A slight adjustment to LARS implements LASSO and calculates all possible LASSO estimates for a given problem.
2. Another variation of LARS efficiently executes forward stagewise linear regression.
3. A rough estimate of the degrees of freedom of a LARS estimate is available, allowing for a calculated prediction error estimate based on C_p . This enables a deliberate choice among the possible LARS estimates.

3.3. Hilbert–Schmidt Independence Criterion LASSO (HSIC-LASSO)

In HSIC-LASSO, a kernelized LASSO is utilized to identify nonlinear relationships between inputs and outputs [28]. By selecting appropriate kernel functions, features that have a strong statistical relationship with the target can be identified using the Hilbert–Schmidt independence criterion, a kernel-based measure of independence. These selected features are not redundant, and the globally optimal solution can be efficiently calculated, making this method suitable for high-dimensional problems.

3.4. Conditional Covariance Minimization (CCM)

The CCM method utilizes kernel-based independence measures to identify a subset of covariates that possess the maximum predictiveness of the response [29]. It achieves feature selection through an optimization problem that involves the trace of the conditional covariance operator.

3.5. Binary Coyote Optimization Algorithm (BCOA)

The binary constrained optimization algorithm (BCOA) [29] is an extension of the constrained optimization algorithm (COA) [30]. It performs feature selection by evaluating the performance of a binary classification algorithm. This is achieved by using the hyperbolic transfer function as a wrapper model to determine the optimal features.

4. Proposed Method

The proposed feature selection method employs the divide-and-conquer technique. This approach is recognized for its recursive execution of the same algorithm at lower levels, and eventually concludes after a finite period of time, as detailed in studies by Rosler [31] and Smith [32].

The divide-and-conquer method has the advantage of being able to tackle large problems efficiently by breaking them down into smaller sub-problems that can be solved individually [19]. This approach enables us to apply feature selection methods on large datasets, even if they are not capable of handling a huge number of features. The process involves dividing the features into smaller sets and then applying the selection methods to these subsets, leading to efficient resolution of the problem.

The aim is to select the top n features from a huge set of features (F). We achieve this by first dividing F into various subsets and finding the top features of these subsets

using a feature selection method (such as the filter method). These subsets are then ranked according to their importance, and a new set of features is selected from this sorted set of features. The new set of top-selected features is further divided into subsets, and the process is repeated until we find the top n features of F . Now the problem we face is that any feature selection method can only select a certain number of features in desired constant time due to its computational efficiency. This limits the usability of that particular feature selection method. However, we can overcome this problem by using the divide-and-conquer approach.

To obtain the set of subsets of features shown in Figure 2, we first decide the size of this set of subsets of features. All the subsets will be of equal size because some of the features will be left out; these features will initially not be included in these subsets but will be evaluated later in the next iteration. L represents this set of features. We remove L from G and divide the remaining features into smaller subsets. A total of X subsets is formed. Each of these subsets is denoted by G_i , where i ranges from 0 to X .

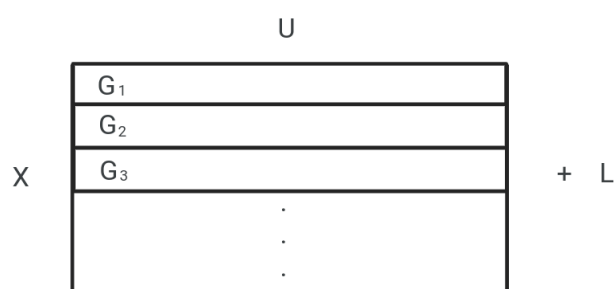


Figure 2. Visual representation of dividing G into subsets.

Now we will apply a filter method for feature selection in these X subsets and score each of the features in these subsets [33]. The filter method used in this paper is based on mutual information. Information gain, also known as mutual information, evaluates the impact of a feature on the accuracy of predicting the target. The concept assumes that if two random variables are independent, the information gain would be zero.

The selection of fuzzy functions is based upon a fuzzy entropy measurement. Since fuzzy entropy may better differentiate model distribution, it is used to evaluate the separability of each characteristic. The membership function used here is the Gaussian function. The Gaussian function is also known as the normal distribution function; it is often used in probability and statistics. The Gaussian function is used to represent fuzzy sets where the degree of membership for a given value is defined by the curve's height at that value. The mean parameter (μ) determines the center of the curve, and the standard deviation parameter (σ) determines the curve's width. The area under the Gaussian curve equals 1, meaning that the membership function is normalized.

Intuitively, the lower the fuzzy entropy of a characteristic, the higher the capacity for discernment of the characteristic. The Shannon probabilistic entropy may be defined as [34]

$$H_1(A) = - \sum_{j=1}^n (\mu_A(x_j) \log \mu_A(x_j) + (1 - \mu_A(x_j)) \log (1 - \mu_A(x_j))) \quad (6)$$

where $\mu_A(x_j)$ are the fuzzy values. This fuzzy entropy measurement is considered a measurement of fuzziness and assesses the overall deviations from the standard series type, i.e., any crisp set A_0 lead to $h(A_0) = 0$. Note that the fuzzy set A with $\mu_A(x_j) = 0.5$ acts as the maximum element in the defined order by H . Newer fuzzy entropy measures [35] may improve the performance.

Apply a filter method for every G_i and arrange the features in descending order of importance in that subset. The first feature of each subset will be the most important feature of that subset. Using this, we will arrange every subset G_i according to its importance, i.e., arrange the subsets G_i in descending order of importance of their first feature. All the

features are now sorted in each subset G_i , and all the subsets G_i are also sorted with respect to each other.

Let,

U = number of features; efficiently selected by feature selection method from a group of features.

U will be the number of features in each subset, any feature selection method would have to select a maximum of U features from each subset.

To divide the features into equal subsets,

Let,

G = some set of features.

Initially, we will take the original set of features,

$G = F$.

Let,

X = number of subsets,

L = set of features left after the division of G into equal subsets,

G_i = subset of G , $i = 1, 2, 3, \dots, X$;

Then,

$\text{sizeof}(G) = X \times U + \text{sizeof}(L)$

$\text{sizeof}()$ represents the size of a set in single dimension.

Given the values of G and U , we can find the value of X and L by,

$$X = \left\lfloor \frac{\text{sizeof}(G)}{U} \right\rfloor,$$

$$\text{sizeof}(L) = \text{sizeof}(G) - X \times U$$

The resulting matrix will have the most important features in the upper left corner, and the least important ones in the bottom right corner (according to the filter method's score). We will split this matrix into two parts according to n (number of features required) and select the features in the upper part of the matrix to form a new set of features. G' will now become the newly obtained set of features, and the process will be repeated until the top n features are selected.

The matrix is divided into two parts, i.e., G_i and G_{i-1} . Each j th feature in G_i has more importance than $(j - 1)$ th. Similarly, for the feature in G_{i-1} ($j = 1, 2, 3, \dots$), we do not know whether if the least important feature of G_i is more important than the most important feature of G_{i-1} or not. This is because the features of a subset are scored with respect to each other (i.e., scoring of features is done within the subset) and not with the features of another subset. Therefore the features that could potentially be the top n features are as follows.

top n features of G_1 ,

or top $n-1$ features of G_1 + top feature of G_2 ,

or top $n-2$ features of G_1 + top 2 features of G_2 ,

or top $n-2$ features of G_1 + top feature of G_2 + top feature of G_3 ,

or top $n-3$ features of G_1 + top 3 features of G_2 ,

or top $n-3$ features of G_1 + top 2 features of G_2 + top feature of G_3 ,

.

.

.

Keeping this in mind, the new subset of features(G) will be selected as,
 $G' = \{\}$, initially it will be an empty set,
 For every subset of features $G_i, i = 1, 2, \dots, X$
 Let,
 $y = \max(0, n - i + 1),$
 $G' = G' + \text{top } y \text{ feature of } G_i$
 The previous set of features G will be updated as,
 $G = G' + L$

This process is repeated until the size of G is smaller than U . Now we apply the filter method to this final set of features to obtain the top n features of the original set of features, as the number of features to select (n) is smaller than U . The final number of features to be selected (i.e., n) should be smaller than U in order for the method to work.

The proposed feature selection method is shown in Figure 3. It was applied to each dataset, with 20, 30, 40, 50, and 60 features selected in each run. The selected features were then used to fit and train the models, which included hyperparameter tuning. The time utilized by the proposed feature selection method was also recorded. The results were then compared to other feature selection methods.

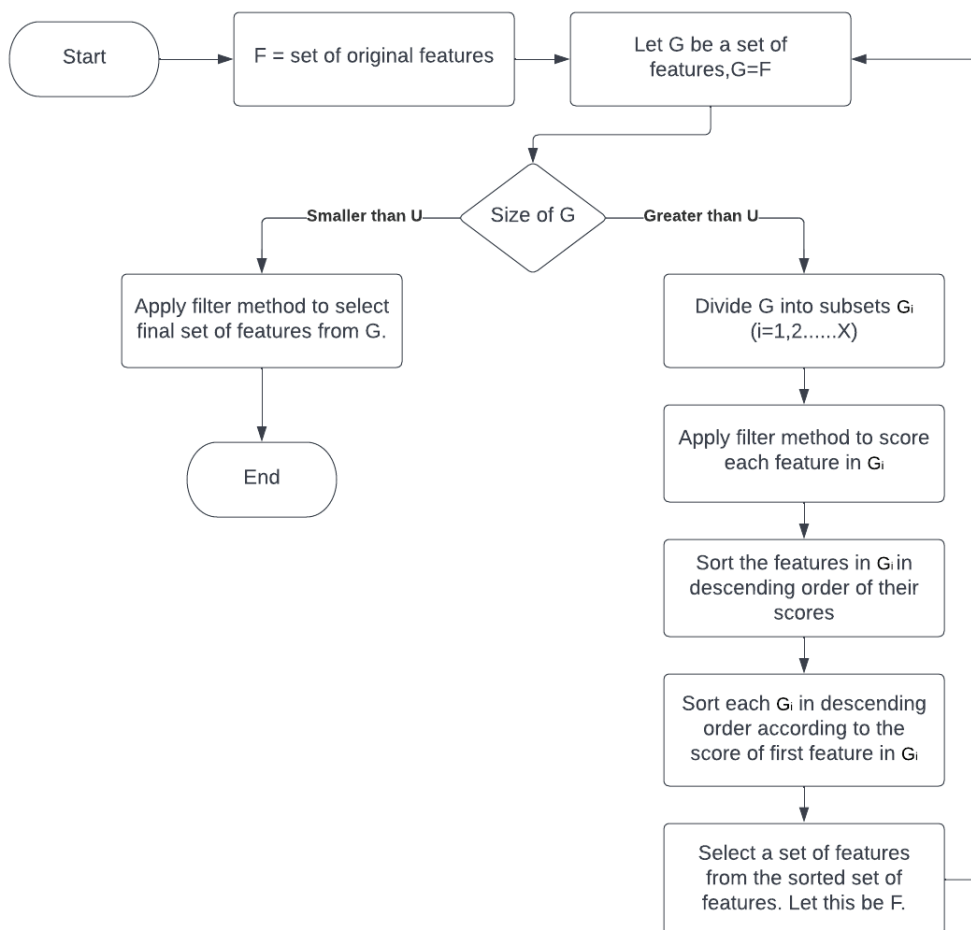


Figure 3. Flowchart of the process.

The pseudocode of the proposed approach is explained in Algorithm 1.

Algorithm 1 Pseudocode of the proposed approach.**Require:**

- import required libraries
- set upper limit of algorithm as 500
- set number of features to extract

Function-1: Data(InputFile, numberOfFeatures)

- define n as numberOfFeatures
- define dataframe by reading input csv file
- apply data pre-processing and split the data into train and test set

Function-2: FeatureDivider(features, upper_limit)

- #This function will divide the given set of features into smaller set of features
- Calculate the number of features that will be in each smaller set of features and divide the bigger set of features accordingly
- Return the set of features and the features left out of the perfect sets

Function-3: Feature_selector(featureSets, k, mod)

- #This function will select top k features from each subset and sort them
- #define the feature selection method as SelectKBest
- for i in every feature subset
 - define curr_df as a feature subset
 - sort the features in this subset
- sort the subsets with themselves by compare in the best features of each subset
- for i in range(number of subsets)
 - for j in range(number of subsets)
 - sort the subsets in ascending order according to their top feature
- return the sorted subset of features as sorted_top_k

Function-4: upper_matrix(sorted_top_k, left, left_features)

- #This function will select features from upper matrix and add left out features to int
- Take out the features from upper matrix of sorted_top_k and append them into final_features
- Append the left out features to final features
- return final_features

Main Function: main_fun(features, upper_limit, mod, k)

- #This function will remove undesired features in each iteration
- while number of features left > upper_limit
 - left_features, feature_sets = feature_divider(features, upper_limit)
 - sorted_top_k = feature_selector(feature_sets, k, mod)
 - final_features = upper_matrix(sorted_top_k, left_features)
 - features = final_features
- run one final iteration which will select out desired number of features
- selector = SelectKBest(mod, k=k)
- selector_fit = selector.fit(x_train[features], y_train)
- top_n_features = selector_fit.get_feature_names_out()
- write these features into a text file

Call the Main Function:

- # call the data function to execute the program
- data(InputFile, 20)

5. Experimental Results

The proposed method will be evaluated against seven leading feature selection techniques: mRMR [26], LARS [27], HSIC-LASSO [28], Fast-OSFS and Scalable [36], group-SAOLA [37], CCM [29], and BCOA [38]. The implementation of the proposed method has been made available in the form of code [39].

5.1. Dataset Used

Datasets were selected from the NCBI database [40]. NCBI is a provider of online biomedical and genomic information resources. The data found in the database is in soft file format and was transformed into CSV using an R program [41,42]. The program also cleaned the data, resulting in a reduction of features. The standardized data was then transformed using a standard scaler to enhance classifier performance.

Table 1 describe the dataset used in the experiment. A total of eight datasets were used for the experiment. The target features of these datasets were divided into two, three, and four classes, with approximately 192 samples in each dataset.

Table 1. Description of the dataset used for the experiment.

Datasets	Samples	Original Features	Cleaned Features	Labels
GDS-1615	127	22,200	13600	Three
GDS-2546	167	12,600	9500	Four
GDS-968	171	12,600	9100	Four
GDS-2545	171	12,600	9300	Four
GDS-3929	183	24,500	19,300	Two
GDS-1962	180	54,600	29,100	Four
GDS-531	173	12,600	9300	Two
GDS-2547	164	12,600	9300	Four

5.2. Classifier Used

The prediction models used are support vector machine (SVM) and random forest [43]. SVM is a classification algorithm that identifies the most influential cases, called support vectors, to form a decision boundary or hyperplane. Random forest, on the other hand, operates in four stages. It selects random samples from the dataset, builds a decision tree for each sample, obtains a prediction from each tree, and chooses the prediction with the most votes through a voting process.

The accuracy of the SVM [44] and random forest classifiers [45] can be improved by tuning their parameters. The experiment utilized grid search cross-validation to optimize the parameters of both classifiers. This approach trains the model using every combination of the specified parameters to find the optimal set, which results in increased accuracy. The parameters of grid search CV for RF used in the experiment were:

```
'n_estimators': np.arange(50,200,30)
'max_features': np.arange(0.1, 1, 0.1)
'max_depth': [3, 5, 7, 9, 50, 100]
'max_samples': [0.3, 0.5, 0.8, 1]
```

where,

- `n_estimators`: The quantity of trees in the forest.
- `max_features`: The number of features considered to find the optimal split.
- `max_depth`: The highest level of the tree.
- `max_samples`: The number of samples taken from X to train each base estimator.
- `np.arange`: Produces evenly spaced values within a specified range.

The parameters of grid search CV for SVM used in the experiment were:

```
'c': [0.01, 1, 5, 10, 100]
'kernel': ('linear', 'poly', 'rbf', 'sigmoid')
'gamma': ('scale', 'auto')
```

where,

- c : A regularization constant.
- kernel: Determines the type of kernel used in the algorithm.
- γ : The kernel coefficient for 'rbf', 'poly', and 'sigmoid'.

5.3. Result and Discussion

Table 2 shows the quantity of selected features, the accuracy obtained, and the duration in seconds of each run of the proposed method for each dataset.

Table 2. Comparative performance of accuracy using random forest and support vector machine.

Datasets	Features Selected	RF	SVM	Time Taken (sec)
GDS1615	20	86.31	83.15	41.29
	30	88.42	86.31	41.2
	40	86.31	82.1	43.25
	50	87.36	88.42	43.87
	60	86.31	88.42	44.97
GDS968	20	76.52	77.41	35.9
	30	74.89	75.84	36.43
	40	73.35	74.21	37.26
	50	76.55	78.8	38.4
	60	78.06	81.26	38.82
GDS531	20	84.49	78.27	29.44
	30	85.29	79.84	22.32
	40	86.86	86.06	22.52
	50	86.09	83.75	22.9
	60	86.83	83.75	23.75
GDS2545	20	72.67	68.76	37.98
	30	78.09	71.04	38.49
	40	75.01	73.29	38.9
	50	70.15	69.53	39.84
	60	75.01	71.81	40.82
GDS1962	20	79.25	72.59	129.13
	30	78.51	73.33	116.72
	40	80	77.77	129.6
	50	79.25	76.29	130.06
	60	79.25	76.29	130.82
GDS3929	20	74.44	69.33	46.38
	30	72.93	69.41	45.37
	40	76.66	67.88	45.91
	50	74.36	67.91	46.37
	60	75.26	67.88	48.54
GDS2546	20	95.2	84	37.87
	30	95.2	85.59	38.01
	40	96	80.8	38.66
	50	95.2	82.4	36.88
	60	95.2	80	38.39

Table 2. Cont.

Datasets	Features Selected	RF	SVM	Time Taken (sec)
GDS2547	20	69.23	67.59	37.77
	30	70.03	68.39	38.2
	40	70.06	67.56	39.1
	50	70.83	68.39	38.07
	60	70.86	65.93	38.79

Table 3 presents the average number of selected features and average classification accuracy obtained over 10 separate runs using SVM and RF on the described datasets. The accuracy of the proposed method is based on the average of five runs, with the number of selected features set to 20, 30, 40, 50, and 60, respectively, and the parameter U fixed at 500 in each run.

Fuzzy backward feature elimination helps to identify the most relevant features in a dataset. By eliminating features that are less important or irrelevant, the algorithm can improve the performance of the model and make it more interpretable. Additionally, it can also help to identify potential sources of noise or bias in the dataset, which can improve the overall accuracy and generalizability of the model.

The proposed features selection method achieved better accuracy than the existing methods in the majority of the cases in a considerably smaller amount of time as compared to the existing methods. It also had considerably low computation requirements, which is beneficial, as this method can be used by anyone on low-performance machines, such as personal computers. This gain in efficiency might not be very observable in small datasets, but it drastically reduces the time required in feature selection in huge dimension datasets.

Figure 4 compares different feature selection techniques using a support vector machine and random forest on different datasets. The techniques compared include traditional methods, such as mutual information and correlation-based feature selection, as well as more recent methods, such as recursive feature elimination and LASSO. The results indicate that the proposed feature selection method surpasses all other techniques in terms of accuracy and computational efficiency. This is likely because the proposed method considers the interactions between features, whereas traditional methods focus solely on individual feature importance. The proposed method's ability to select the most relevant features for the task at hand results in a more robust and accurate model.

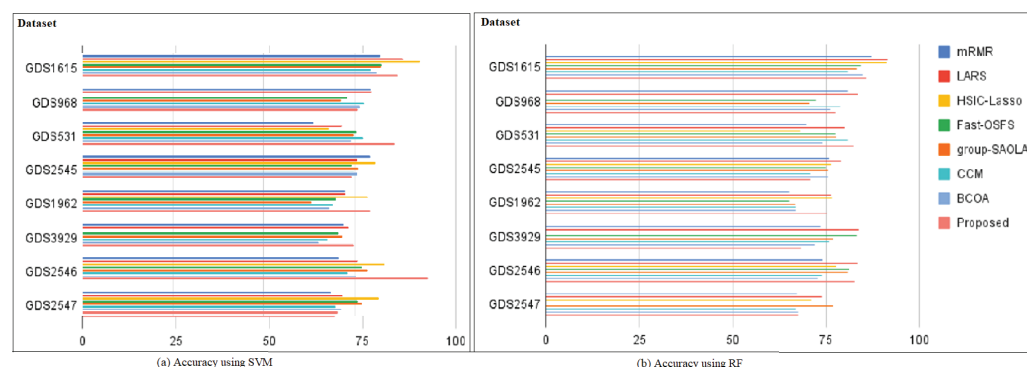


Figure 4. Comparison of different feature selection techniques using support vector machine and random forest on different datasets.

The empty spaces in Table 3 mean that those specific feature selection methods do not run on the datasets. The accuracy achieved by the proposed method is better than other methods in most cases for SVM and random forest classifiers. The number of features selected was higher than in other methods. As the number of features selected was increased,

the accuracy was increased in some cases but did not vary much, which shows that the method is rational.

This method does not require high computational power and can run on most machines. This was possible because of the divide-and-conquer approach, which reduced the complexity of the problem. The time taken by the proposed method was significantly less, keeping the required time under 1 m for five datasets, 30 s for one dataset, and about 2 m for one dataset.

Table 3. Comparative performance of different feature selection techniques over five runs for each method. Format: accuracy (number of features).

Classifier	Datasets	mRMR	LARS	HSIC LASSO	Fast OSFS	Group SAOLA	CCM CCM	BCOA BCOA	Proposed Work
SVM	GDS1615	87.37(40)	91.67(26)	91.35(18)	84.31(17)	83.13(12)	80.82(29)	84.9(33)	88.42(50)
	GDS968	80.87(39)	83.73(38)	-	72.41(19)	70.53(14)	78.82(34)	76.19(32)	81.26(60)
	GDS531	69.78(30)	79.96(27)	67.93(4)	77.43(26)	77.7(11)	80.82(30)	74.17(32)	86.06(40)
	GDS2545	75.9(34)	79.02(33)	76.4(33)	74.95(18)	75.55(12)	70.82(30)	75.4(29)	73.29(40)
	GDS1962	65.12(39)	76.56(32)	76.81(31)	65.15(24)	66.59(10)	66.82(40)	66.89(35)	77.77(40)
	GDS3929	73.57(41)	83.78(41)	-	83.11(40)	76.97(21)	75.82(39)	72.12(41)	69.41(30)
	GDS2546	74.13(33)	83.51(32)	77.69(27)	81.25(26)	80.88(17)	73.82(35)	72.98(32)	85.59(30)
	GDS2547	67.31(39)	73.88(32)	71.16(12)	73.13(23)	76.85(24)	66.82(28)	67.35(26)	68.39(30)
RF	GDS1615	81.96(32)	88.24(20)	92.88(22)	82.34(15)	82.26(13)	79.55(31)	81.08(30)	88.42(30)
	GDS968	79.44(44)	79.77(42)	-	72.84(19)	71.28(18)	77.53(41)	76.42(40)	78.06(60)
	GDS531	63.69(23)	71.44(20)	67.82(4)	75.48(14)	74.67(16)	77.36(23)	73.92(21)	86.86(40)
	GDS2545	79.31(31)	75.81(33)	80.64(33)	74.16(14)	76.05(12)	74.82(34)	75.63(33)	78.09(30)
	GDS1962	72.37(29)	72.41(30)	78.45(42)	69.88(21)	63.28(13)	69.17(32)	67.95(30)	80(40)
	GDS3929	71.94(29)	73.44(28)	-	70.49(28)	71.56(15)	67.5(28)	65.13(24)	76.66(40)
	GDS2546	70.53(36)	75.86(34)	83.09(45)	77.04(25)	78.46(18)	72.9(36)	75.28(31)	96(40)
	GDS2547	68.44(22)	71.68(24)	81.67(32)	75.85(30)	77.1(20)	69.7(25)	71.28(24)	70.86(60)

6. Application of the Proposed Work

The proposed method, apart from genomic datasets, can be used in various areas, such as in healthcare datasets, where the number of features for a given individual can be massive (i.e., blood pressure, blood group, resting heart rate, height, weight, immune system status, surgery history, nutrition, and existing conditions), in financial datasets, where the number of features for a given stock can be quite large, and in ecological datasets.

Contour shape analysis is a demonstration of data analysis in infinite dimensions, specifically, analysis on the projective spaces of a Hilbert space. This method involves using high-dimensional approximations and operates within the framework of a Hilbert manifold.

The proposed method uses the divide-and-conquer technique with fuzzy backward feature elimination (FBFE) that helps to find the important features quickly and accurately. To show the robustness of the proposed method, it is applied to eight different datasets taken from the NCBI database. The only requirement is that the selection method scores each feature for the divide-and-conquer approach to work. This opens up a vast area of research where few feature selection methods can be used over this method.

7. Conclusions

This paper presents a feature selection method to select the top n important features from a huge dimensional dataset. The novelty of our method is that it can run on a huge set of features in less time and use less computational power. It uses a divide-and-conquer approach to select the top n important features from the datasets. The proposed method performed well compared to other state-of-the-art feature selection methods on the datasets.

This method can be made more accurate by using other feature selection methods, such as wrapper and exhaustive feature selection methods, instead of filter methods.

In this paper, the focus was only on genomics datasets, as these are widely accepted, and only a few dimension reduction algorithms work well on these datasets. In future work, we aim to revise our feature selection method and use wrapper methods with this proposed method.

Author Contributions: Conceptualization—P.S.R., W.A. and M.D.A.; Methodology—A.T. and H.S.; Validation—W.A., P.S.R., H.S. and M.D.A.; Writing—P.S.R., W.A., M.D.A., A.T. and H.S.; Visualization—M.D.A. and A.T.; Funding acquisition—W.A. and M.D.A. All authors have read and agreed to the published version of the manuscript.

Funding: This research work was funded by Institutional Fund Projects under grant no. (IFPIP: 1335-611-1443). The authors gratefully acknowledge the technical and financial support provided by the Ministry of Education and King Abdulaziz University, DSR, Jeddah, Saudi Arabia.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: Not applicable.

Conflicts of Interest: The authors declare no conflict of interest.

References

- Chandrashekar, G.; Sahin, F. A survey on feature selection methods. *Comput. Electr. Eng.* **2014**, *40*, 16–28. [CrossRef]
- Altman, N.; Krzywinski, M. The curse(s) of dimensionality. *Nat. Methods* **2018**, *15*, 399–400. [CrossRef] [PubMed]
- Hughes, G. On the mean accuracy of statistical pattern recognizers. *IEEE Trans. Inf. Theory* **1968**, *14*, 55–63. [CrossRef]
- Baştanlar, Y.; Özuysal, M. Introduction to machine learning. In *miRNomics: MicroRNA Biology and Computational Analysis*; Humana Press: Totowa, NJ, USA, 2014; pp. 105–128.
- Law, M.H.; Figueiredo, M.A.; Jain, A.K. Simultaneous feature selection and clustering using mixture models. *IEEE Trans. Pattern Anal. Mach. Intell.* **2004**, *26*, 1154–1166. [CrossRef] [PubMed]
- Kuncheva, L.I.; Matthews, C.E.; Arnaiz-González, Á.; Rodríguez, J.J. Feature selection from high-dimensional data with very low sample size: A cautionary tale. *arXiv* **2020**, arXiv:2008.12025.
- Kohavi, R.; John, G.H. Wrappers for feature subset selection. *Artif. Intell.* **1997**, *97*, 273–324. [CrossRef]
- Blum, A.L.; Langley, P. Selection of relevant features and examples in machine learning. *Artif. Intell.* **1997**, *97*, 245–271. [CrossRef]
- Jović, A.; Brkić, K.; Bogunović, N. A review of feature selection methods with applications. In Proceedings of the 2015 38th International Convention on Information and Communication Technology, Electronics and Microelectronics (MIPRO), Opatija, Croatia, 25–29 May 2015; pp. 1200–1205.
- Liu, H.; Sun, J.; Liu, L.; Zhang, H. Feature selection with dynamic mutual information. *Pattern Recognit.* **2009**, *42*, 1330–1339. [CrossRef]
- Battiti, R. Using mutual information for selecting features in supervised neural net learning. *IEEE Trans. Neural Netw.* **1994**, *5*, 537–550. [CrossRef]
- Wah, Y.B.; Ibrahim, N.; Hamid, H.A.; Abdul-Rahman, S.; Fong, S. Feature Selection Methods: Case of Filter and Wrapper Approaches for Maximising Classification Accuracy. *Pertanika J. Sci. Technol.* **2018**, *26*, 329–340.
- El Aboudi, N.; Benhlila, L. Review on wrapper feature selection approaches. In Proceedings of the 2016 International Conference on Engineering & MIS (ICEMIS), Agadir, Morocco, 22–24 September 2016; pp. 1–5.
- Bolón-Canedo, V.; Sánchez-Marono, N.; Alonso-Betanzos, A.; Benítez, J.M.; Herrera, F. A review of microarray datasets and applied feature selection methods. *Inf. Sci.* **2014**, *282*, 111–135. [CrossRef]
- Liu, H.; Zhou, M.; Liu, Q. An embedded feature selection method for imbalanced data classification. *IEEE/CAA J. Autom. Sin.* **2019**, *6*, 703–715. [CrossRef]
- Liu, H.; Motoda, H. *Feature Extraction, Construction and Selection: A Data Mining Perspective*; Springer Science & Business Media: Berlin/Heidelberg, Germany, 1998; Volume 453.
- Aziz, R.; Verma, C.; Srivastava, N. A fuzzy based feature selection from independent component subspace for machine learning classification of microarray data. *Genom. Data* **2016**, *8*, 4–15. [CrossRef]
- Yasmin, G.; Das, A.K.; Nayak, J.; Pelusi, D.; Ding, W. Graph based feature selection investigating boundary region of rough set for language identification. *Expert Syst. Appl.* **2020**, *158*, 113575. [CrossRef]
- Reimann, M.; Doerner, K.; Hartl, R.F. D-ants: Savings based ants divide and conquer the vehicle routing problem. *Comput. Oper. Res.* **2004**, *31*, 563–591. [CrossRef]
- Song, Q.; Ni, J.; Wang, G. A fast clustering-based feature subset selection algorithm for high-dimensional data. *IEEE Trans. Knowl. Data Eng.* **2011**, *25*, 1–14. [CrossRef]

21. Zhao, Q.; Meng, D.; Xu, Z. A recursive divide-and-conquer approach for sparse principal component analysis. *arXiv* **2012**, arXiv:1211.7219.
22. Sun, Y.; Todorovic, S.; Goodison, S. Local-learning-based feature selection for high-dimensional data analysis. *IEEE Trans. Pattern Anal. Mach. Intell.* **2009**, *32*, 1610–1626.
23. Armanfard, N.; Reilly, J.P.; Komeili, M. Local feature selection for data classification. *IEEE Trans. Pattern Anal. Mach. Intell.* **2015**, *38*, 1217–1227. [CrossRef]
24. Tibshirani, R. Regression shrinkage and selection via the lasso. *J. R. Stat. Soc. Ser. (Methodol.)* **1996**, *58*, 267–288. [CrossRef]
25. Chen, X.; Yuan, G.; Nie, F.; Huang, J.Z. Semi-supervised Feature Selection via Rescaled Linear Regression. In Proceedings of the IJCAI, Melbourne, Australia, 19–25 August 2017; Volume 2017, pp. 1525–1531.
26. Peng, H.; Long, F.; Ding, C. Feature selection based on mutual information criteria of max-dependency, max-relevance, and min-redundancy. *IEEE Trans. Pattern Anal. Mach. Intell.* **2005**, *27*, 1226–1238. [CrossRef] [PubMed]
27. Efron, B.; Hastie, T.; Johnstone, I.; Tibshirani, R. Least angle regression. *Ann. Stat.* **2004**, *32*, 407–499. [CrossRef]
28. Yamada, M.; Jitkrittum, W.; Sigal, L.; Xing, E.P.; Sugiyama, M. High-dimensional feature selection by feature-wise kernelized lasso. *Neural Comput.* **2014**, *26*, 185–207. [CrossRef] [PubMed]
29. Chen, J.; Stern, M.; Wainwright, M.J.; Jordan, M.I. Kernel feature selection via conditional covariance minimization. *Adv. Neural Inf. Process. Syst.* **2017**, *30*, 6949–6958.
30. Pierezan, J.; Coelho, L.D.S. Coyote optimization algorithm: A new metaheuristic for global optimization problems. In Proceedings of the 2018 IEEE congress on evolutionary computation (CEC), Brisbane, Australia, 10–15 June 2018; pp. 1–8.
31. Rösler, U. On the analysis of stochastic divide and conquer algorithms. *Algorithmica* **2001**, *29*, 238–261. [CrossRef]
32. Smith, D.R. The design of divide and conquer algorithms. *Sci. Comput. Program.* **1985**, *5*, 37–58. [CrossRef]
33. Guyon, I.; Gunn, S.; Nikravesh, M. *Feature Extraction: Studies in Fuzziness and Soft Computing*; Springer Science & Business Media: Berlin/Heidelberg, Germany, 2006.
34. Luukka, P. Feature selection using fuzzy entropy measures with similarity classifier. *Expert Syst. Appl.* **2011**, *38*, 4600–4607. [CrossRef]
35. Parkash, O.; Gandhi, C. Applications of trigonometric measures of fuzzy entropy to geometry. *Int. J. Math. Comput. Sci.* **2010**, *6*, 76–79.
36. Yu, L.; Liu, H. Feature selection for high-dimensional data: A fast correlation-based filter solution. In Proceedings of the 20th International Conference on Machine Learning (ICML-03), Washington, DC, USA, 21–24 August 2003; pp. 856–863.
37. Yu, K.; Ding, W.; Wu, X. LOFS: A library of online streaming feature selection. *Knowl.-Based Syst.* **2016**, *113*, 1–3. [CrossRef]
38. de Souza, R.C.T.; de Macedo, C.A.; dos Santos Coelho, L.; Pierezan, J.; Mariani, V.C. Binary coyote optimization algorithm for feature selection. *Pattern Recognit.* **2020**, *107*, 107470. [CrossRef]
39. Available online: <https://github.com/git-arihant/feature-importance> (accessed on 25 December 2022).
40. Available online: <https://www.ncbi.nlm.nih.gov/> (accessed on 25 December 2022).
41. Available online: <https://github.com/jranaraki/NCBIdataPrep> (accessed on 25 December 2022).
42. Afshar, M.; Usefi, H. High-dimensional feature selection for genomic datasets. *Knowl.-Based Syst.* **2020**, *206*, 106370. [CrossRef]
43. Sheykhoumousa, M.; Mahdianpari, M.; Ghanbari, H.; Mohammadimanesh, F.; Ghamisi, P.; Homayouni, S. Support vector machine versus random forest for remote sensing image classification: A meta-analysis and systematic review. *IEEE J. Sel. Top. Appl. Earth Obs. Remote. Sens.* **2020**, *13*, 6308–6325. [CrossRef]
44. Yuanyuan, S.; Yongming, W.; Lili, G.; Zhongsong, M.; Shan, J. The comparison of optimizing SVM by GA and grid search. In Proceedings of the 2017 13th IEEE International Conference on Electronic Measurement & Instruments (ICEMI), Yangzhou, China, 20–22 October 2017; pp. 354–360.
45. Sumathi, B. Grid search tuning of hyperparameters in random forest classifier for customer feedback sentiment prediction. *Int. J. Adv. Comput. Sci. Appl.* **2020**, *11*, 173–178.

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

Article

Fuzzy Rule Based Adaptive Block Compressive Sensing for WSN Application

Dibyalekha Nayak ¹, Kananbala Ray ¹, Tejaswini Kar ¹ and Sachi Nandan Mohanty ^{2,*}

¹ School of Electronics Engineering, KIIT Deemed to be University, Bhubaneswar 751024, Odisha, India

² School of Computer Science & Engineering, VIT-AP University, Amaravati 522237, Andhra Pradesh, India

* Correspondence: sachinandan.m@vitap.ac.in

Abstract: Transmission of high volume of data in a restricted wireless sensor network (WSN) has come up as a challenge due to high-energy consumption and larger bandwidth requirement. To address the issues of high-energy consumption and efficient data transmission adaptive block compressive sensing (ABCS) is one of the optimum solution. ABCS framework is well capable to adapt the sampling rate depending on the block's features information that offers higher sampling rate for less compressible blocks and lower sampling rate for more compressible blocks. In this paper, we have proposed a novel fuzzy rule based adaptive compressive sensing approach by leveraging the saliency and the edge features of the image making the sampling rate selection completely automatic. Adaptivity of the block sampling ratio has been decided based on the fuzzy logic system (FLS) by considering two important features i.e., edge and saliency information. The proposed framework is experimented on standard dataset, Kodak data set, CCTV images and the Set5 data set images. It achieved an average PSNR of 34.26 and 33.2 and an average SSIM of 0.87 and 0.865 for standard images and CCTV images respectively. Again for high resolution Kodak data set and Set 5 dataset images, it achieved an average PSNR of 32.95 and 31.72 and SSIM of 0.832 and 0.8 respectively. The experiments and the result analysis show that proposed method is efficacious than the state of the art methods in both subjective and objective evaluation metrics.

Keywords: block compressive sensing (BCS); fuzzy decision; saliency detection; edge detection

MSC: 94A08; 03B52

1. Introduction

The wireless sensor network have been successfully applied to different fields, such as disaster management, medical investigation, military battlefield, surveillance and home automation [1–3]. However, the WSN lags in handling huge amount of image data due to the bandwidth and energy limitation as higher the complexity of the process greater the energy requirement [1,3]. To avoid the large data transmission efficient compression is very much essential which will be suitable for WSN application. Compressive sensing (CS) based compression is most widely used compression technique in the WSN [4–7]. In contrast to the scalar data such as temperature, soil moisture etc. image data put much more burden in the WSN transmission and also poses the challenges due to their restricted resources. Therefore, the compression of the raw image is highly required to save the energy and space. Different image compression algorithms, which are based on the transform based approach and follow the Nyquist sampling theorems (i.e., JPEG [8], JPEG2000 [9]), are not applicable for the restricted WSN application. JPEG [8] and JPEG2000 [9] are not preferable in the restricted WSN because of its coding complexity, power consumption and channel errors. The Compressive sensing (CS) [10] is a decade old technique, which doesn't follow the Nyquist sampling theorem and is suitable for both data acquisition and compression. The CS provides sampling and compression as a related process rather than

the two separate process. The study of CS revealed that the signal can still be recovered even from a lower number of samples. At the time of transmission data loss occur in WSN due to the multi path processing. However with CS the complete data can be recovered from the insufficient data with minimal loss [4–6]. This sampled data recovery system of CS and advantages of adaptive BCS have motivated us to design an automatic and adaptive block image compression algorithm based on CS for WSNs.

Since CS based approach deals with vectored representation of the image, it becomes difficult to handle the complete image at a time. The solution is to use block-based approach of compressive sensing through BCS [11] instead of processing the complete image in one go. Due to block-based processing, the size of the sampling matrix is also small that lead to lower computational complexity. Here the sampling matrix size is fixed for all the blocks. Since, the image blocks are varying according to their feature content, applying a fixed sampling ratio for all the blocks is not an optimum choice. Considering this, Adaptive Block Compressive Sensing (ABCS) [12] technique is developed where the blocks with higher feature content block should be compressed at lower rate than the blocks with lower feature content to maintain the clarity of the reconstructed image and to restore the data loss. In the ABCS method, the selection of the blocks was done by using a threshold value depending on the feature content of the block.

2. Related Work

This section describes the literature review related to different CS based approaches for image compression. Monika, et al. [12] proposed a review on ABCS method which gives the detailed description regarding the challenges and different adaptivity consideration for block selection. Zha, et al. [13] proposed Discrete Cosine Transform (DCT) based sparsity computation followed by sparsity dependent ABCS method. Here the sampling rate output will be any of the four predetermined levels. Therefore, it is not completely automatic image adaptive sampling ratio detection method. Monika, et al. [14] proposed coefficient permuted ABCS method suitable for under water WSN using fewer samples. Xu, et al. [15] proposed an adaptive perceptual BCS scheme which combines sparsity and perceptual sensibility to decide the sampling rate allocation of the blocks. Heng, et al. [16] proposed a fuzzy logic based ABCS method by considering the standard deviation and sparsity of the image block. The authors have also developed a fuzzy based recovery algorithm that maintains the quality of the reconstruction. This method is efficient in terms of quality of the reconstruction. However the authors are silent about the complexity analysis. Sum, et al. [17] proposed texture feature based approach for adaptivity of the blocks which may not be suitable for all varieties of images. Many saliency based the ABCS approaches [18–20] can be found in the literature, where the authors used saliency information for deciding adaptivity of the blocks. Different saliency detection approaches can be found in the literature [21–23] which are used for wild life monitoring.

In the literature authors have applied different features such as entropy [24], wavelet co-efficient [25], standard deviation [26] to find the adaptive sampling ratio of the blocks suitable for gray scale images. These methods were successful in maintaining the quality of the reconstructed image. Zhao, et al. [27] proposed gradient based ABCS method However it suffers from blocking artifacts in the reconstructed image. Zhao, et al. [28] proposed multi shaped block strategies based adaptive block sampling ratio for real time application on IoT. Due to the use of different block sizes, it increases the burden on the processor. Canh, et al. [29] proposed an edge preserving CS recovery technique by considering the nonlocal and the histogram information of the image in the form of gradient. The method doesn't have much significant improvement over existing methods [12]. Li, et al. [30] proposed an adaptive block sampling rate detection method by considering the error between the blocks. However, reconstructed image has poor quality. Monika, et al. [31] proposed an energy adaptive block selection based ABCS method. As this method has less energy consumption, hence suitably applied for underwater based IoT images. Gambhir, et al. [32] proposed the edge and fuzzy transform based image compression using the

standard deviation and the saliency information of the blocks. Wang, et al. [33] proposed an adaptive rate compressive sensing method by considering the statistical data analysis for the sparsity calculation. The performance of the method is better, but statistical data analysis increases the execution time and power [34]. Kazemi, et al. [35] proposed a multi focus based image fusion technique by using the adaptive sampling rate. The adaptivity of the block is decided by using textural information. Multi focus fusion approach of images increases the complexity of the method [36].

Most of the methods found in the literature have considered only single feature based approach for the adaptive sampling measurement and only few have considered dual feature based approach, some have experimented on gray images, some have the limitation of higher energy consumption and some have the limitation of higher delay and power requirements [34]. Moreover, most of the methods are non-automatic sampling allocation schemes except the fuzzy based approach.

Therefore, motivated by the state of the requirements and advantages of fuzzy based approach we have designed a fuzzy based ABCS model. Towards this goal we have proposed a novel fuzzy rule based approach for finding the adaptive sampling value selection by using Mamdani fuzzy rule approach. Because of which the sampling rate allocation is completely automatic. Since the proposed work considers two features, it can be better capture the image features for finding the adaptive sampling value.

The contribution of the paper are as follows:

- A novel fuzzy rule based compressive sensing technique for WSN application.
- An automatic threshold generation technique for adaptive sampling by considering two important features.
- Probability based saliency map detection
- Simple provision for handling oversampling cases.

The remainder of this paper are arranged as follows. Section 3 focusses on the proposed method in detail, Section 4 deals with the experiments and result analysis and finally Section 5 concludes the paper.

3. Proposed Method

In this section we have described about proposed fuzzy logic based adaptive sampling rate generation system by considering two important features i.e., edge and saliency information of the block). The proposed method is divided into 4 different stages. The different stages are as in Section 3.1. Block selection and Feature Extraction, Section 3.2. Adaptive sampling phase, Section 3.3. CS measurement phase, Section 3.4. Reconstruction phase. The architecture diagram of the proposed work flow is shown in the Figure 1. The detailed description of the individual stages are described in the following subsection.

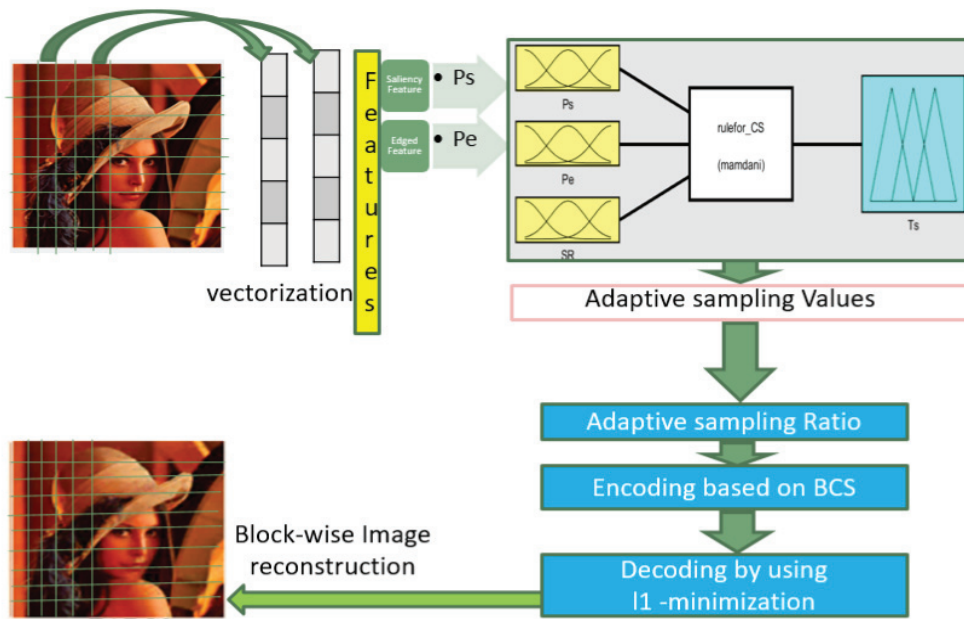


Figure 1. The Proposed Model.

3.1. Block Selection and Feature Extraction

Literature review revealed that the block-based processing for measuring matrix selection has the advantage of less execution time and lower memory requirement in contrast to the traditional CS. However, selection of block size is the key component in BCS approach. The size of the block not only decides the sampling rate and the measuring matrix but also improves the reconstructed image quality. Again, assigning the fixed measuring matrix or the sampling ratio to all the blocks without consideration of different feature content is not an optimal solution. Some image blocks may contain more crucial information while others may not. Therefore, to achieve an effective compression, adaptive sampling rate allocation is highly preferable.

With the aim of maintaining the quality of reconstruction, we have considered two main features i.e., edge and the saliency of the blocks. In this work, the adaptiveness of sampling ratio is decided based on the combination of the edge features and the salient feature data of the block. In this subsection, we primarily discuss about block size selection and suitable feature extraction. Neither a very small size nor a very large size of the block is suitable for optimum feature extraction. Here, the image is divided into 16×16 or 32×32 , equal sized non overlapping blocks for feature extraction.

3.1.1. Saliency Detection

The image has been considered as a combination of the different saliency information. Here the saliency map detection is based on the high probability pixels of an image. In this method we have considered the histogram of the image, which gives idea about the probability of occurrence of the image pixels. From the probability information based on histogram data the high probability pixels are obtained as follows.

The threshold value t_1 detection in (1) is based on the highest probability value of an image pixel $I(x, y)$, which gives an ideal bench mark for the high probability pixel occurrence criteria.

$$t_1 = \frac{\max(\text{Probability of all pixel})}{2} \quad (1)$$

$$y_1(x_1) : \text{Set of high probability pixels denoted as} \quad (2)$$

$$y_1(x_1) \in \{P(x, y) \geq t_1\}$$

$P(x, y)$: denotes Probability of a pixel $I(x, y)$ at location (x, y) ,

Contrast values of the high probability pixels, which are obtained by using the global contrast method are calculated by using (3)

$$D(x1, x1 \in N_1) = \sum_{x,y=1}^{i,j} abs(y1(x1) - I(x, y)) \quad (3)$$

where, $I(x, y)$ is the image pixel value at location (x, y) , N_1 is the total number of pixels in $y1$.

The threshold value is obtained based on the sum of average of all contrast values depending upon the number of high probability pixel values in an image. The saliency can be detected based on the threshold as given in (4)

$$Ts = \frac{\sum_{i=1}^{N_1} D(x1_i)}{N_1} \quad (4)$$

The saliency map ($Smap(x, y)$) is obtained using (5)

$$Smap(x, y) = \begin{cases} 1 & D(x1) > Ts \\ 0 & otherwise \end{cases} \quad (5)$$

By counting the number of 1's the block wise saliency percentage (P_s) is obtained as given in (6)

$$P_s(\%) = \frac{\text{Total count of 1}}{B \times B} \times 100 \quad (6)$$

where, B is size of image block.

3.1.2. Edge Detection

The edge map of the image is obtained by Canny Edge detection algorithm [37]. Canny edge detection algorithm proceeds through a number of stages to get wide range of edges in an optimal way. The first stage of the algorithm is primarily for smoothening of images, as it reduces the noise level by using the Gaussian filter. In the second stage, the intensity of gradient of the pixel values and the edge thinning technique has been applied to prevent from maximum suppression. Edge map is obtained by double thresholding and connectivity analysis. The edge map is divided into non-overlapping blocks of size $B \times B$. In each block the percentage of 1's denoted as P_e is evaluated as given in (7)

$$P_e(\%) = \frac{\text{Total count of 1}}{B \times B} \times 100 \quad (7)$$

The proposed method developed an automatic sampling rate allocation by considering dual feature based fuzzy logic.

For the visual reference saliency map and the edge map of Lena image are shown in Figure 2. The evaluation of saliency feature (P_s) and edge feature (P_e) of the block was done from the saliency map and the edge map as shown in Figure 1. After evaluation of the (P_s) and (P_e) the adaptive sampling of the block was carried out by fuzzy rule based approach, using fuzzy inferencing system described in next subsection.

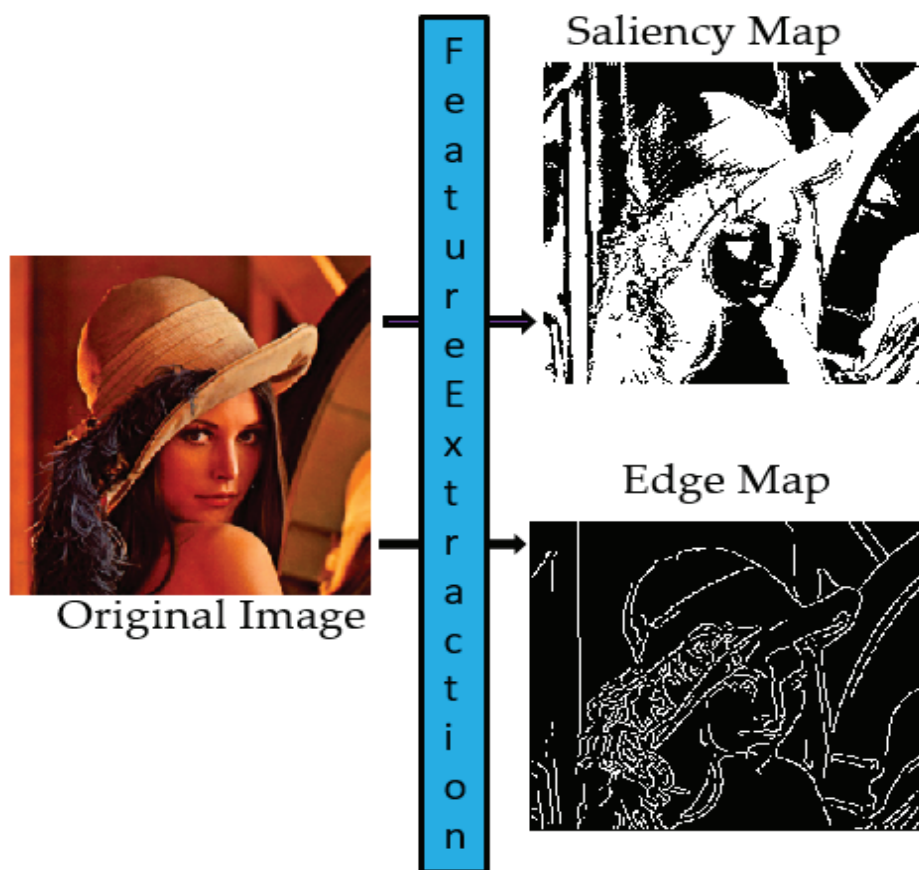


Figure 2. Saliency and Edge Map Representation.

3.2. Adaptive Sampling Phase

In this phase, a feature based adaptive allocation of the sampling rate has been carried out. The adaptive sampling rate allocation was done using fuzzy logic system. This phase is having different decision sub-stages such as designing of fuzzy logic system, adaptive sampling rate allocation by using fuzzy rules, and over sampling correction.

3.2.1. Fuzzy Logic System

L. Zadeh has developed fuzzy logic system, which is a convenient technique in improving the decision making in resource constraint networks i.e., WSN due to its conservation of the resources and operative performance. FLS provide much intelligent solution in the regulator problems by imitating the human thought process. As shown below the FLS has four main components; those are fuzzifier, an inference engine, fuzzy rules, and a defuzzifier. In the fuzzification process, the fuzzifier converts the craps data into fuzzy sets (linguistic terms are low, medium, high etc.) by considering a membership function. The membership function map the non-fuzzy input into the fuzzy linguistic terms and vice-versa. This process enables the knowledge of the linguistic terms. So many different memberships functions have been developed and deployed in the literature include triangular, trapezoidal, Gaussian, piecewise linear and singleton function. The choice of the membership function has done based on their experience and assessment of the researcher. The membership functions are useful in both the fuzzification and defuzzification process. Fuzzification is followed by fuzzy inferencing process via fuzzy if-then rule base with condition and conclusions. FLS gives the defuzzified membership function as output. Defuzzification process gives the crisp output.

3.2.2. Adaptive Sampling Rate Allocation by Using Fuzzy Rules

In this work a fuzzy rule-based approach has been applied for deciding adaptive sampling of the blocks in an automatic way. The primary motivation to design fuzzy rule based approach for deciding adaptive sampling value is to make the process fully automatic. In this regard, the authors took the advantage of the feature content of each block and the sampling ratio (SR) to decide the adaptive sampling value of the blocks via fuzzy rule-based approach. The feature content of the blocks are captured by the saliency percentage (P_s) and the edge percentage (P_e) of the blocks as described in Sections 3.1.1 and 3.1.2 respectively. Our proposed FLS has three inputs and one output as illustrated in Figure 3. The fuzzy rules are created by carrying out number of experiments on different images. The fuzzy rules have been framed and put in a tabular form in Table 1. While designing the fuzzy rules, care has been taken to maintain the minimum reconstruction quality of the blocks by applying a fixed sampling value. Therefore, irrespective of the feature content of the blocks, a minimum amount of sampling ratio has been assigned to avoid the blocking artifacts. The details of the mamdani based Fuzzy rules are given in Table 1 using various linguistic variables for all inputs and outputs. The triangular and trapezoidal membership functions are studied for representation of the linguistic variables for input and output and the best one is selected. Various possible combination of rules have been exercised and 65 rules have generated and out of those rule 40 basic rules are shown in Table 1. The range of the values of the each and every graph is decided after carrying number of tests upon the image dataset. The proposed fuzzy rule-based approach is well capable to reduce the computation burden by using FLS design. The output value of the FLS is decided by considering the saliency, edge percentage and the sampling ratio.

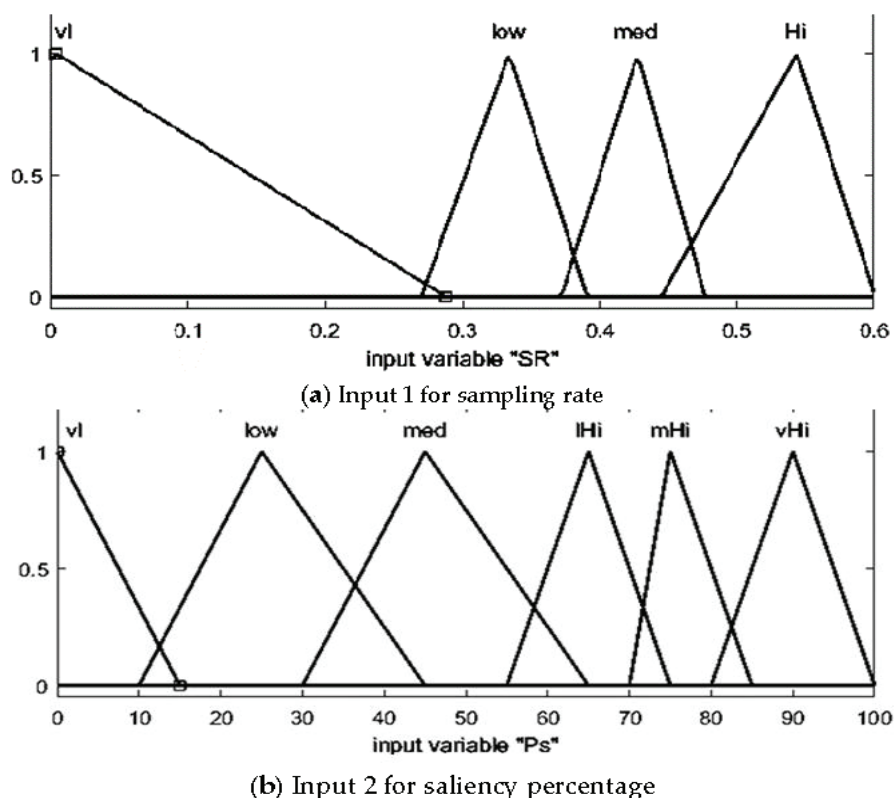
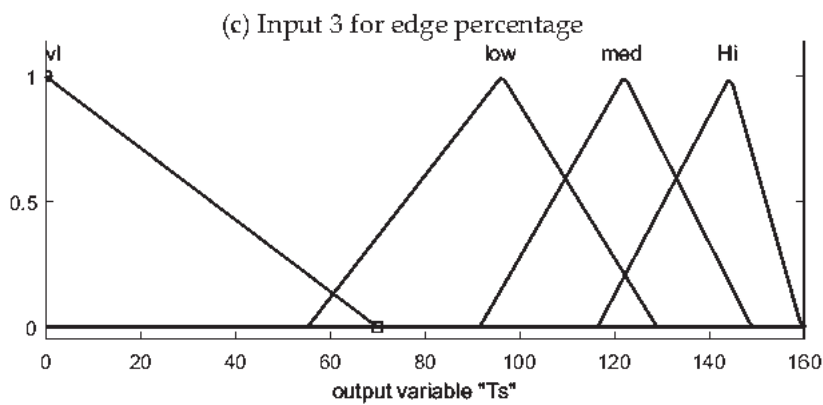
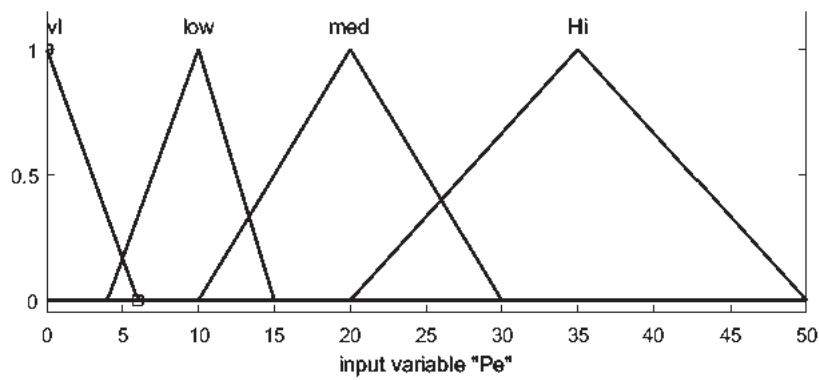


Figure 3. Cont.



(d) Output for adaptive sampling: threshold

Figure 3. Graphical representation of Input and Output membership function.

Table 1. Illustration of the Fuzzy Rules.

Sampling Rate	Fuzzy Rules		
	Saliency Percentage (P_s)	Edge Percentage (P_e)	Adaptive Sampling Values (Th)
Hi	Vl	Vl	Low
Hi	Vl	Med	Low
Hi	Vl	Low	Low
Hi	Vl	Hi	Med
Hi	Low	Low	Med
Hi	Low	Hi	Med
Hi	Low	Med	Med
Hi	Low	Vl	Med
Hi	Hi-Med	Med	Med
Hi	Hi-Med	Low	Med
Hi	Hi-Med	Hi	Hi
Hi	Hi-Med	Vl	Med
Hi	Low-Med	Hi	Med
Hi	Low-Med	Low	Med
Hi	Med-Hi	Vl	Med
Hi	Med-Hi	Med	Hi
Hi	Med-Hi	Low	Med
Med	Vl	Vl	Vl
Med	Vl	Med	Vl
Med	Vl	Low	Vl
Med	Vl	Hi	Low

Table 1. Cont.

Sampling Rate	Fuzzy Rules		Adaptive Sampling Values (Th)
	Saliency Percentage (P_s)	Edge Percentage (P_e)	
Med	Low	Low	Low
Med	Low	Hi	Med
Med	Low	Med	Low
Med	Low	VI	Low
Med	Hi-Med	Med	Low
Med	Hi-Med	Low	Low
Med	Hi-Med	Hi	Med
Med	Hi-Med	VI	Low
Med	Low-Med	Hi	Med
Med	Low-Med	Low	Low
Low	Hi-Med	VI	Low
Low	Low-Med	Hi	Low
Low	Low-Med	Low	Low
Low	Med-Hi	VI	Low
Low	Med-Hi	Med	Low
Low	Med-Hi	Low	Low
Low	Med-Hi	Hi	Low
Low	Low-Hi	VI	VI
Low	Low-Hi	Low	VI
Low	Low-Hi	Hi	Low

VI = Very low, Med = Medium, Hi = High, Med-Hi = Medium High, Low-Med = Low Medium, Hi-Med = High medium.

3.2.3. Over Sampling Corrections

The output of the FLS is optimum. However, in some cases the adaptive sampling value exceeds the block size, where the block is processed as an uncompressed block. This condition is known as oversampling of the blocks. Oversampling can lead to lower performance of the compression algorithm therefore it must be avoided. To avoid this condition, the oversampling of the blocks is checked after generation of the output of the FLS. In case the block suffers from oversampling a corrective measure, need to be taken. To deal with the oversampling, first the oversampled value is obtained by subtracting the block size from the adaptive measured sampling data. Subsequently the oversampled values are distributed among all the non-over sampled blocks as per an algorithm discussed below.

Let B_s : denotes the block size, A_m denotes adaptively measured sampling value, N_b : denotes total number of blocks, A_s : denotes adaptive sampling, O_s denotes over sampling value and n denotes as the number of non over sampling blocks. Using these notations the algorithm to handle the oversampling of the blocks is given as follows in Algorithm 1.

Algorithm 1: Over Sampling Corrections

```

For  $i = 1 : N_b$ 
   $O_s = 0; n = 0;$ 
  If ( $A_m(i) \geq B_s$ )
    # over sampled block
    Then  $O_s(i) = A_m(i) - B_s;$ 
     $A_s(i) = A_m(i) - O_s(i);$ 
     $O_s = O_s(i) + O_s;$ 
     $n = n + 1;$ 
  Else
    # Non oversampled block
     $A_s(i) = A_m(i) + O_s/n;$ 
  End if
End for

```

3.3. Compressive Sensing Measurement Phase

In the proposed approach, the adaptive sampling is followed by compressive sensing for the encoding and decoding of the image data. In this section a detailed description of the CS and BCS have been given. According to conventional CS theory an image can be reconstruct from few linear measurements if the signal is sparse enough. A signal is sparse in the transform domain if most of the elements are zero and nearly sparse if prominent feature portions are zero or nearly zero. The basic idea in the CS is that a signal can be recovered from very less decoded information. Suppose, x represents a signal of length L and K sparse in the transform domain, ψ represents the transform basis of x of size $L \times L$. θ gives the transform domain representation of the x having only K coefficients. x is represented as $x = \theta\psi$. The linear transform process of getting the value of y (size $M < L$) from x is known as the CS sampling. The linear equation illustrating CS sampling is given in (8).

$$y_i = \Phi x = \Phi \theta \psi \quad (8)$$

Here, Φ : represents random measuring matrix of size $M \times L$. Then x can be recovered completely by considering $M = O(K \log \frac{L}{K})$ measurements and by solving the convex optimization problem. To make the signal reconstruction easier x represented in the l1 norm form as in (11).

$$x = \underset{x: y = \Phi x}{\operatorname{argmin}} \|x\|_1 \quad (9)$$

The 2D image signal has converted into the 1D signal, followed by conversion into sparse signal by using transforms such as DCT [Discrete Cosine Transform] or DWT [Discrete wavelet Transform]. Further, the l1-minimization process for the 1-D signal reconstruction.

Block Compressive Sensing Method for Image Compression

The CS based approach for compressive sensing is very impressive in terms of performance. However, the major issue is pertaining to the reconstruction of the large sized image with lower computational complexity. Moreover, handling the entire image data, requires large size measuring matrix or sampling matrix which leads to undesirable increase in the complexity and larger storage requirement. To mitigate these issues BCS scheme has been proposed.

In BCS, image is divided into $B \times B$ non-overlapping blocks and each block is then considered as a 2D signal. A fixed sampling rate r is applied to the image block x_i and the sampled value acquired denoted as y_i are related as given in (10).

$$y_i = \Phi^{(B)} x_i = \Phi^{(B)} \psi^{(B)} X_i \quad (10)$$

Here the $\Phi^{(B)} = n_B \times B^2$ represents the measuring matrix of the image block, which is constructed by selecting $n_B = (r \times B^2)$, rows from a matrix of size $(B^2 \times B^2)$ and the transform matrix $\psi^{(B)}$ of size $(B^2 \times B^2)$. Both the transform matrix and the sparse vector will form $x_i = \psi^{(B)} X_i$. The computation of sampling matrix for the image block is space and time saving process than the entire image sampling matrix calculation. The encoded signal is transmitted through the transmitter by having all the adaptive sampling information.

3.4. Reconstruction Phase [10]

In this phase the decoding of the compressed data is carried out. For the image reconstruction many different techniques can be found in the literature. However, l1-minimization based technique for the image reconstruction is preferable owing to its higher accuracy [4]. The BCS-SPL approach has attracted attention of many researchers due to incomplete data handling and lower reconstruction artifacts, which is suitable for real time application. In this work, we have built a fuzzy rule based adaptive sampling design approach for BCS-SPL algorithm.

4. Experimental Design and Result Analysis

To evaluate the proposed method we have experimented on four different types of dataset. In the first phase we have considered 6 different.tiff standard colour images i.e., Lena, Zelda, Couple, House, Baboon, Peppers image. In the second phase, we have considered Kodak dataset [38] comprising of 25 colour images. In the third phase, we have considered the low resolution real life CCTV images [39] collected from the National park surveillance system. To evaluate the performance of the proposed method for different resolution of the images here we have considered some high resolution images from the Set5 dataset [40]. All the images considered above are resized to a size of 256×256 . To evaluate the performance of the method both subjective and objective evaluation has been performed. The subjective analysis has been done based on the perceptual view whereas, objective analysis has been done based on suitable performance parameters evaluation such as PSNR (Peak Signal to Noise Ratio) and Structural similarity Index matrix measurement (SSIM) [41]. All simulations are carried out considering non overlapping blocks of size 16×16 . The objective evaluation of the proposed method has been given in Table 2, based on PSNR [13] and SSIM [41] comparison of different state of the art methods such as BCS-SPL [11], ENT-ABCS [6], STD_BCS-SPL [26] and EABCS [31] with proposed FABCS approach. Although different reconstruction approaches can be found in the literature [10,12,31], but l1-minimization based reconstruction approach has reduced blocking artifacts with superior reconstruction image quality [10]. Therefore, to do a fair comparison l1 minimization based reconstruction approach is applied for all the methods.

The subjective evaluation is made based on the visual quality of the reconstructed image. For all the considered methods the reconstructed images of standard dataset are shown in Figure 4. Figure 4a illustrates the original images of Lena, Zelda, Couple, Baboon and House. Figure 4b–f indicate the reconstructed images using BCS-SPL, STD-BCS-SPL, ENT-ABCS, EABCS and FABCS at a SR of 0.5. From Figure 4 it is found that the perceptual quality of the proposed FABCS method is the best among all the approaches considered for evaluation. Out of all the methods considered for evaluation, the BCS-SPL [9] method is based on a fixed compression rate for all the blocks of image. Which is having a restriction to enhance the reconstruction quality based on block feature content. The STD-BCS-SPL [26] framework considered the standard deviation as a threshold to extract the block wise information content. The higher the information content of the block, the lesser is the sampling ratio and the lower the information content of the block, higher is the sampling ratio. The STD-BCS-SPL method is applicable for restricted WSN application [26]. It has lower computational complexity with better-reconstruction quality of the image. However, Due to the manual calculation the sampling rate allocation is time consuming. In ENT-ABCS [6] method, the entropy of the block was selected to generate the threshold for deciding the adaptive sampling. It relied only on one feature for finding the block adaptivity. In the EABCS method [30] the energy of the block was considered for the adaptive sampling rate selection suitable for IoUT (under water based Internet on things) application. All the state of art methods discussed above have considered only single feature for deciding adaptive sampling. In the proposed approach the dual feature combination can capture even the finer details of the block. The bold faced value in each row of Table 2 indicates the highest value of the performance parameters of corresponding image. It is observed from Table 2 that the proposed method out performed all the other existing methods based on average PSNR and SSIM values for all the images considered for evaluation. Moreover, the fuzzy rule-based evaluation of the adaptive sampling is well capable to provide an automatic image adaptive sampling ratio generation to reduce the computation with better PSNR and SSIM values than the state-of-the-art methods as depicted in the Table 2.



Figure 4. Reconstructed Lena, Zelda, Couple, Baboon and House image by using different ABCS methods (a) original image, (b) BCS-SPL, (c) STD-BCS-SPL, (d) ENT-ABCS, (e) EABCS, (f) FABC at 0.5 sampling rate.

Table 2. Performance comparison of different state of the art methods based on PSNR and SSIM parameters for different state of the art methods at 0.3, 0.4, and 0.5 SR.

Images		BCS-SPL [11]		ENT-ABCS [6]		STD_BCS-SPL [26]		EABCS [31]		FABCS	
		PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM
0.3	Zelda	30	0.751	32.2	0.811	32.3	0.838	32	0.75	33.41	0.8565
	Couple	31	0.742	32.6	0.797	32.8	0.809	32.3	0.77	32.88	0.8345
	House	30	0.735	32.2	0.756	32.5	0.764	32	0.74	32.7	0.8125
	Baboon	25	0.574	28.9	0.563	28.8	0.522	28.9	0.55	30.56	0.6788
	Lena	29	0.812	31.9	0.778	31	0.75	31.3	0.75	33.41	0.8342
	Average	28.9	0.723	31.56	0.741	31.49	0.7366	31.28	0.7132	32.792	0.8031
0.4	Zelda	31	0.861	32.9	0.846	33.4	0.871	32.8	0.8	33.89	0.8641
	Couple	32	0.855	33.4	0.833	33.7	0.8464	33	0.81	34.47	0.8568
	House	32	0.822	32.9	0.784	33.9	0.8441	32.6	0.78	34.34	0.8588
	Baboon	29	0.626	29.3	0.673	30.3	0.6987	29.3	0.57	31.03	0.7837
	Lena	31	0.841	32.6	0.817	33	0.8417	32	0.79	34.03	0.8723
	Average	30.8	0.801	32.24	0.7907	32.8	0.82038	31.94	0.7496	33.552	0.8471
0.5	Zelda	31	0.889	33.5	0.867	34	0.886	33.4	0.82	34.99	0.9029
	Couple	33	0.883	34	0.856	34.4	0.8693	33.4	0.83	35.52	0.9152
	House	32	0.85	33.5	0.807	34.5	0.8711	33.1	0.81	35.11	0.8704
	Baboon	31	0.756	30.5	0.717	31	0.7412	29.6	0.61	31.23	0.8032
	Lena	32	0.881	33.2	0.844	33.6	0.8642	32.7	0.82	34.49	0.8881
	Average	32	0.852	32.93	0.818	33.5	0.846	32.4	0.778	34.27	0.876

The second experiment was done for the real-life low resolution original CCTV images [39]. The evaluation of the FABCS performance using CCTV images are presented in the Table 3 and compared with four state of the art methods. Figure 5 shows the reconstructed CCTV images by different methods. It is observed from Figure 5 that the dual features consideration of the proposed approach is well capable to maintain the reconstruction quality even for the low resolution CCTV images. Dual feature selection scheme ensures better feature extraction and hence the higher PSNR and SSIM values as compared to the state of the art methods. Table 3 represents the comparison of the state-of-the-art methods with proposed FABCS method by using the real life CCTV image from a National Park [36]. The bold faced value in each of Table 3 indicates the highest values of the performance parameters of corresponding image. The Table 3 shows that the proposed method has higher PSNR and SSIM values than the other methods. The SSIM and the PSNR values of the FABCS indicate that it has restored more features during reconstruction phase than other state of the art approaches. The subjective evaluation is made based on the visual quality of the reconstructed image and shown in Figures 4 and 5 respectively.

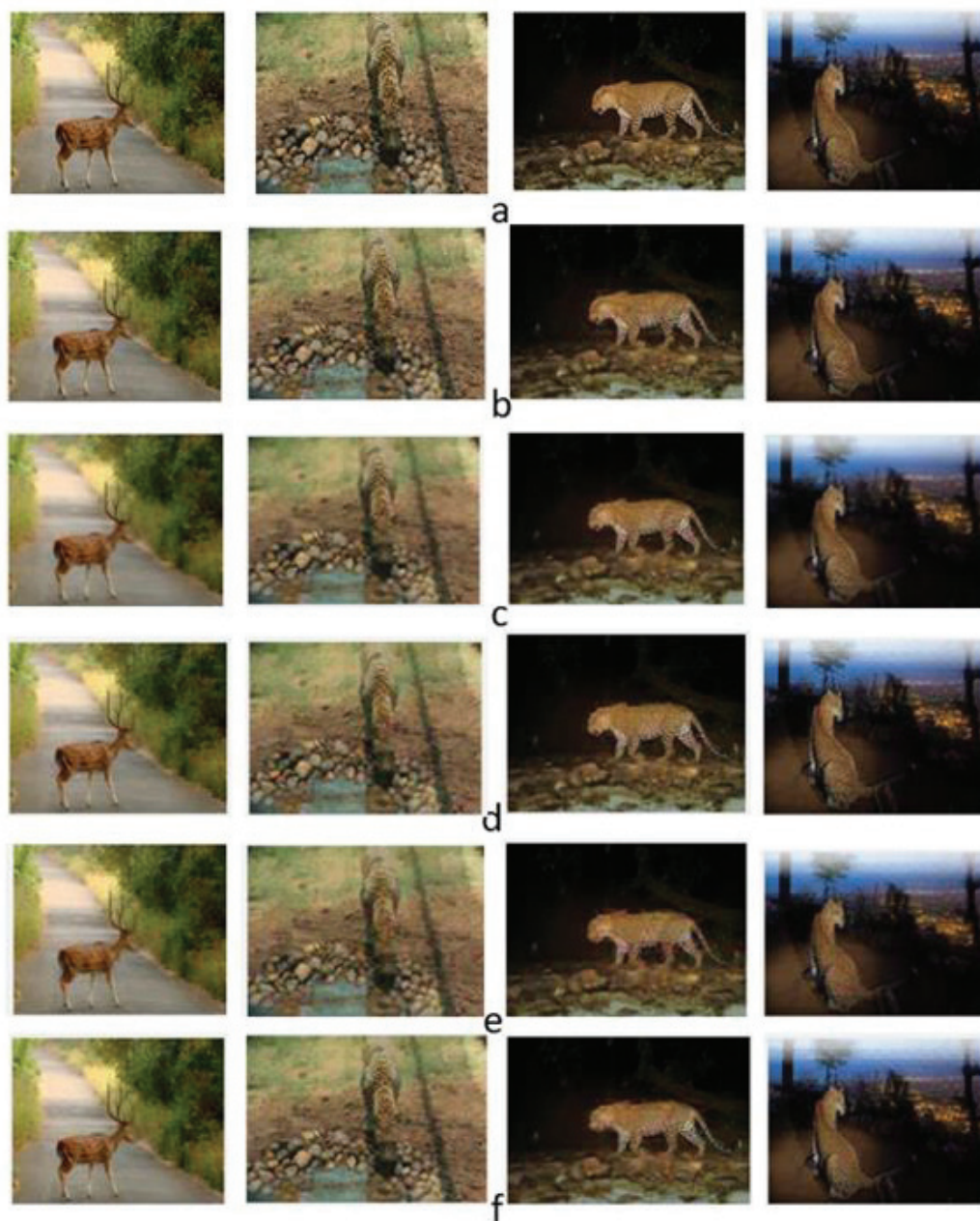


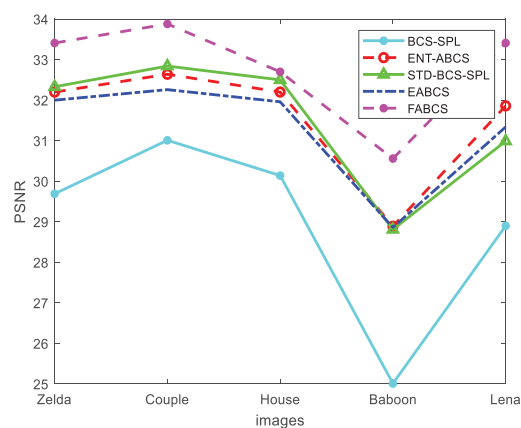
Figure 5. Reconstructed Deer, Leopards images at SR of 0.5 for different methods. (a) Original image, (b) BCS-SPL, (c) STD-BCS-SPL, (d) ENT-ABCS, (e) EABCS, (f) FABCS.

The proposed FABCS method outperformed the other ABCS methods in terms of both subjective and objective evaluation metrics. In the proposed method the adaptive sampling allocation has been done by using the Fuzzy Logic System. Further experiments were carried out to study the effect of variation of SR values on the performance parameter PSNR and SSIM and placed in Figure 6. Figure 6a–e illustrate the statistical representation of average PSNR and SSIM values of the FABCS method with other state of the art methods at SR value of 0.3, 0.4 and 0.5 respectively for standard data set images (i.e., Lena, Zelda, Couple, Baboon and House).

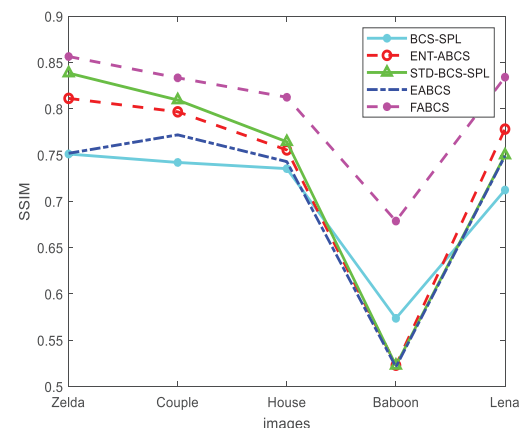
Table 3. Performance comparison of the state of the art methods at 0.5 SR on CCTV images.

Images	BCS-SPL [11]		ENT-ABCS [6]		STD-ABCS [26]		EABCS [31]		FABCS	
	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM	PSNR	SSIM
Deer	31.84	0.82	32	0.85	32.7	0.86	31.99	0.85	33.05	0.9
Leopard1	27.89	0.708	29.63	0.78	29.6	0.789	30.65	0.812	31.92	0.83
Leopard2	31.32	0.821	32.38	0.81	33.23	0.84	30.24	0.828	33.68	0.869
Leopard3	32.13	0.838	31.71	0.83	33.44	0.85	31.23	0.806	33.82	0.878
Leopard4	31.42	0.827	32.44	0.85	32.58	0.842	32.1	0.85	32.73	0.85
Leopard5	30.56	0.81	32.56	0.85	32.82	0.865	29.84	0.795	33.18	0.86
Average	30.86	0.8	32	0.82	32.01	0.84	31.01	0.83	33.06	0.86

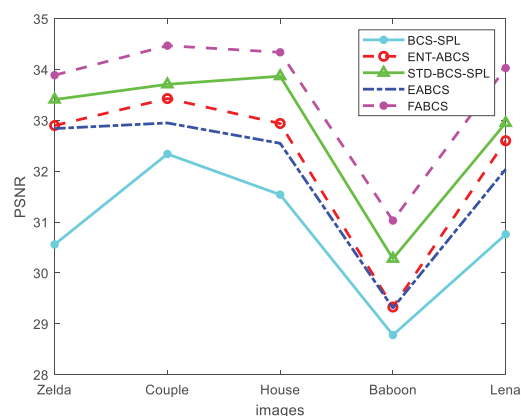
From the image wise analysis it is observed from Figure 6 that the baboon image has the lowest SSIM and PSNR values due to the less sharp features. Couple, House and Lena have SSIM and PSNR values close to each other. The Zelda has the highest SSIM and PSNR values among all the five images as the image has high information content. From the image wise performance analysis, it can be concluded that the proposed method outperformed the state of the art methods. Figure 6g,h illustrate the average PSNR and SSIM comparison of different state of the art methods with FABCS method at 0.3, 0.4 and 0.5 sampling rates for standard dataset images. Figure 6i,j illustrate the average PSNR and SSIM comparison of different state of the art methods with FABCS method at 0.3, 0.4 and 0.5 sampling rates for Kodak dataset.



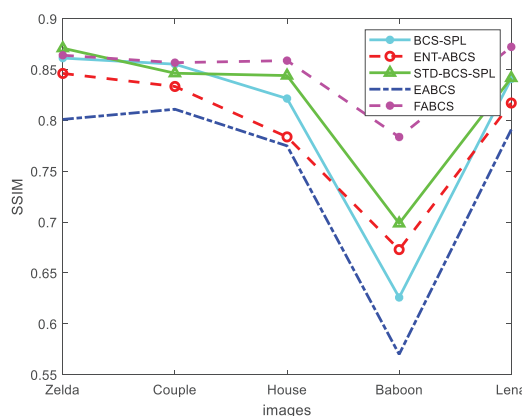
(a) PSNR at 0.3 SR



(b) SSIM at 0.3 SR

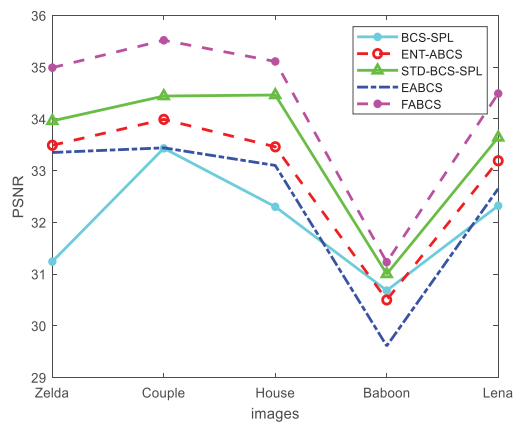


(c) PSNR at 0.4 SR

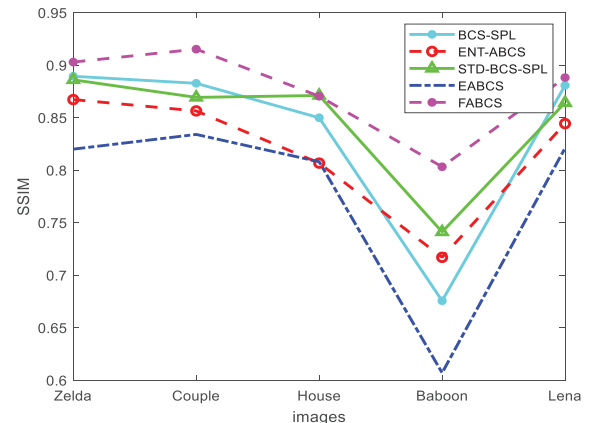


(d) SSIM at 0.4 SR

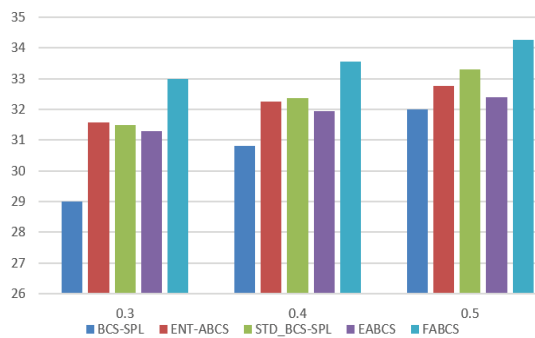
Figure 6. Cont.



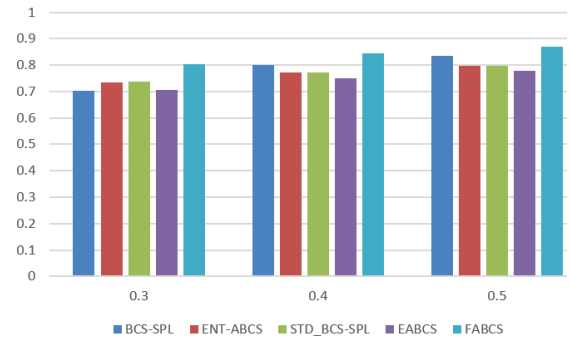
(e) PSNR at 0.5SR



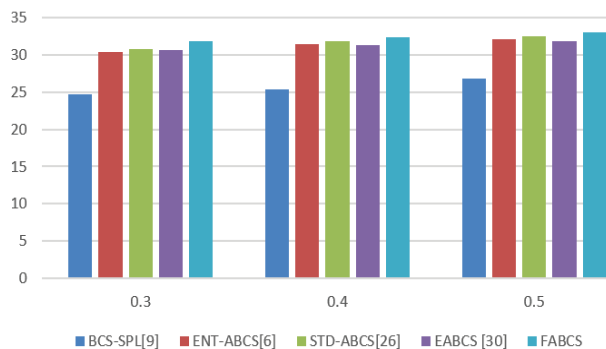
(f) SSIM at 0.5 SR



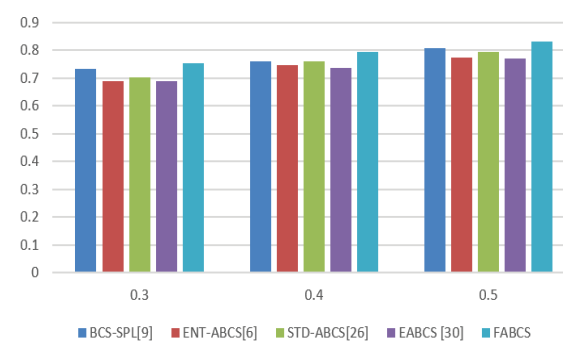
(g) Average PSNR



(h) Average SSIM



(i) Average PSNR for the Kodak dataset



(j) Average SSIM for the Kodak dataset

Figure 6. (a–f) Performance comparison of different algorithms on Standard images and (g,h), (i,j) average PSNR and SSIM for standard images and Kodak dataset images.

To check the performance of the algorithm in the real time image some CCTV images have been taken into consideration. The performance comparison has illustrated in Figure 7. The Figure 7a–f illustrate the detail analysis of the proposed FABCS method over the state of the art methods for CCTV images for different sampling rates at 0.3, 0.4, and 0.5 respectively. The average PSNR and SSIM value are shown in Figure 7g,h which indicate that the proposed FABCS method has the highest PSNR and SSIM than all the state of the art methods by using real time images.

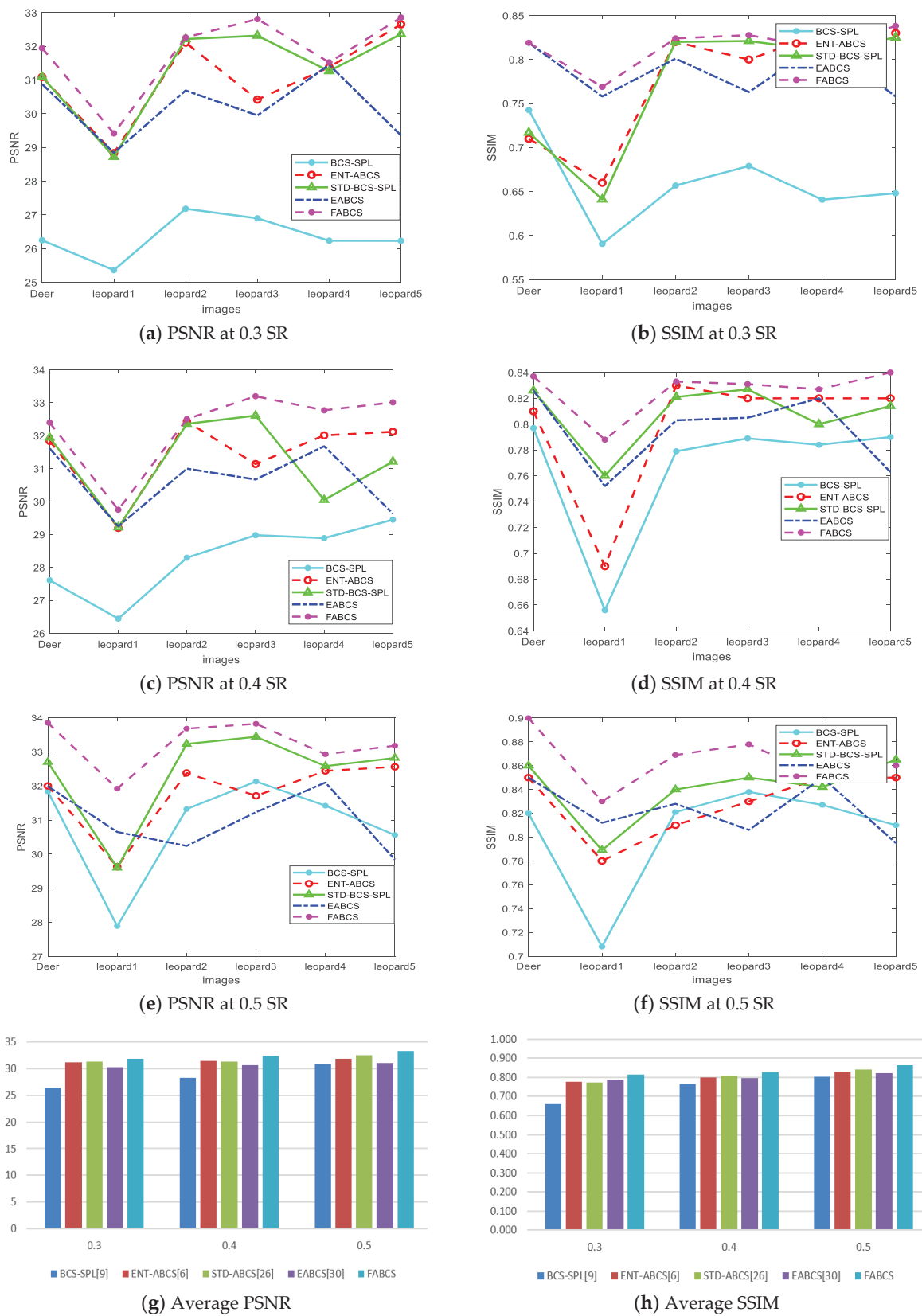


Figure 7. Performance comparison of different algorithms on CCTV images based on PSNR and SSIM.

In order to study the performance of the proposed algorithm on high resolution images we have also experimented on Set5 dataset and the results are placed in Figure 8. It is apparent from Figure 8 that FABCS has also performed well for high resolution images in terms of higher PSNR and SSIM values than the state of the art methods.

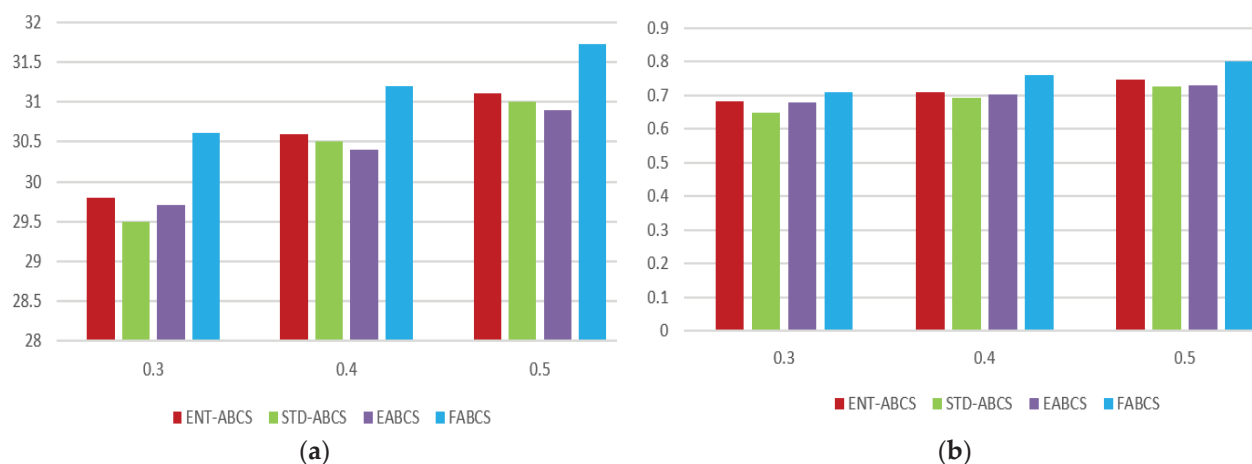


Figure 8. Performance comparison of different algorithms on Set5 Test images based on (a) Average PSNR and (b) Average SSIM.

Figure 9 shows the performance comparison of the fuzzy based adaptive sampling method with the non-fuzzy based adaptive sampling method. For the comparison purpose we have computed the sampling value of the proposed method with and without fuzzy based approach followed by BCS. The performance is evaluated for standard dataset images and the average value is considered. Average PSNR and SSIM value of the fuzzy based ABCS method is 2% higher than the non-fuzzy based ABCS method. The performance of the fuzzy method is higher due to better and a balanced distribution of the sampling rate allocation for the blocks.

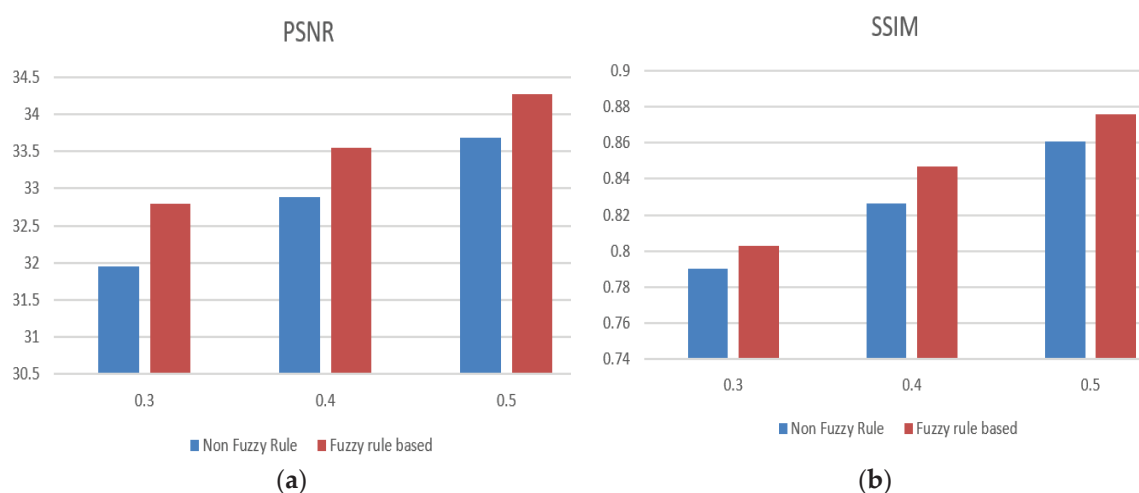


Figure 9. Performance comparison of the Fuzzy and Non fuzzy based approach in terms of (a) PSNR and (b) SSIM.

We have tested our algorithm performance with the 4 different images data set and the statistical analysis of the PSNR and SSIM evaluation has been depicted in the Figures 6–8. Table 4 shows the comparison of average execution time (ET) of different algorithms for standard images of block size (16×16).

It is evident from Table 4 that the proposed FABCS method has less ET than the BCS_SPL for a block size of 16×16 block size. The FABCS has ET lesser than BCS-SPL [11] and slightly more ET than the other state of art methods, due to consideration of dual features and 65 fuzzy rules for automation of the process.

Table 4. Comparison of the ET of different methods in second.

Methods	BCS-SPL [11]	ENT-ABCS [6]	STD_BCS-SPL [26]	EABCS [31]	FABCS
ET (s)	0.376	0.245	0.26	0.252	0.32

Result analysis shows that the proposed FABCS is efficient in terms of better performance measure parameters. The main attraction is the consideration of dual feature selection to capture the content of the images and simple l1 minimization based image reconstruction approach. The fuzzy based block sampling ratio detection has made the algorithm automatic and more effective. Moreover, the proposed approach is followed by a corrective action which ensured that the algorithm should not over sample the block data.

5. Conclusions

The proposed FABCS method is an efficient approach to give a better alternative solution to deal with the practical problem of the data transmission in a wireless sensor network. The proposed method has reduced the computational burden by introducing the fuzzy based adaptive sampling ratio selection in an automatic way. Dual feature selection for threshold detection ensures minimal data loss. The proposed method performed equally well for all variety of test images including standard dataset image, low resolution CCTV image and high resolution Set5 dataset images. The proposed method achieved an average PSNR of 34.26 and average SSIM if 0.87 standard test images. Similarly, it has achieved an average PSNR of 33.2 and SSIM of 0.865 for CCTV images. However, the proposed method can still be improved by considering other features like colour, texture etc. and optimizing the fuzzy rules.

Author Contributions: Conceptualization, D.N. and T.K.; Methodology, D.N. and S.N.M.; Software, D.N. and T.K.; Validation, T.K.; Formal analysis, K.R. and T.K.; Investigation, S.N.M.; Resources, K.R.; Writing—original draft, D.N.; Writing—review & editing, D.N.; Supervision, K.R. and S.N.M.; Funding acquisition, S.N.M. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Data Availability Statement: Not applicable.

Acknowledgments: Article Processing Charge is supported by VIT-AP University Amaravati, Andhra Pradesh INDIA.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Majid, M.; Habib, S.; Javed, A.R.; Rizwan, M.; Srivastava, G.; Gadekallu, T.R.; Lin, J.C.-W. Applications of Wireless Sensor Networks and Internet of Things Frameworks in the Industry Revolution 4.0: A Systematic Literature Review. *Sensors* **2022**, *22*, 2087. [CrossRef]
2. Tavli, B.; Bicakci, K.; Zilan, R.; Barcelo-Ordinas, J.M. A survey of visual sensor network platforms. *Multimedia Tools Appl.* **2011**, *60*, 689–726. [CrossRef]
3. Charalampidis, P.; Fragkiadakis, A.G.; Tragos, E.Z. Rate-Adaptive Compressive Sensing for IoT Applications. In Proceedings of the 2015 IEEE 81st Vehicular Technology Conference (VTC Spring), Glasgow, UK, 11–14 May 2015; pp. 1–5. [CrossRef]
4. Fayed, S.; Youssef, S.M.; El-Helw, A.; Patwary, M.; Moniri, M. Adaptive compressive sensing for target tracking within wireless visual sensor networks-based surveillance applications. *Multimedia Tools Appl.* **2015**, *75*, 6347–6371. [CrossRef]
5. Li, R.; Duan, X.; Li, X.; He, W.; Li, Y. An Energy-Efficient Compressive Image Coding for Green Internet of Things (IoT). *Sensors* **2018**, *18*, 1231. [CrossRef] [PubMed]

6. Rippel, O.; Bourdev, L. Real-time adaptive image compression. In Proceedings of the 34th International Conference on Machine Learning, Sydney, Australia, 6–11 August 2017; Volume 70, pp. 2922–2930.
7. Zammit, J.; Wassell, I.J. Adaptive Block Compressive Sensing: Toward a Real-Time and Low-Complexity Implementation. *IEEE Access* **2020**, *8*, 120999–121013. [CrossRef]
8. Xiao, W.; Wan, N.; Hong, A.; Chen, X. A Fast JPEG Image Compression Algorithm Based on DCT. In Proceedings of the 2020 IEEE International Conference on Smart Cloud (SmartCloud), Washington, DC, USA, 6–8 November 2020; pp. 106–110. [CrossRef]
9. Gungor, M.A.; Gencol, K. Developing a Compression Procedure based on the Wavelet Denoising and Jpeg2000 Compression. *Optik* **2020**, *218*, 164933. [CrossRef]
10. Donoho, D.L. Compressed sensing. *IEEE Trans. Inf. Theory* **2006**, *52*, 1289–1306. [CrossRef]
11. Gan, L. Block compressed sensing of natural images. In Proceedings of the 2007 15th International Conference on Digital Signal Processing, Cardiff, UK, 1–4 July 2007; pp. 403–406.
12. Monika, R.; Samiappan, D.; Kumar, R. Adaptive block compressed sensing—A technological analysis and survey on challenges, innovation directions and applications. *Multimedia Tools Appl.* **2020**, *80*, 4751–4768. [CrossRef]
13. Zha, Z.; Liu, X.; Zhang, X.; Chen, Y.; Tang, L.; Bai, Y.; Wang, Q.; Shang, Z. Compressed sensing image reconstruction via adaptive sparse nonlocal regularization. *Vis. Comput.* **2016**, *34*, 117–137. [CrossRef]
14. Monika, R.; Dhanalakshmi, S.; Kumar, R.; Narayanamoorthi, R. Coefficient Permuted Adaptive Block Compressed Sensing for Camera Enabled Underwater Wireless Sensor Nodes. *IEEE Sens. J.* **2021**, *22*, 776–784. [CrossRef]
15. Xu, J.; Qiao, Y.; Fu, Z. Adaptive Perceptual Block Compressive Sensing for Image Compression. *IEICE Trans. Inf. Syst.* **2016**, *99*, 1702–1706. [CrossRef]
16. Heng, S.; Aintongkham, P.; Vo, V.N.; Nguyen, T.G.; So-In, C. Fuzzy Adaptive-Sampling Block Compressed Sensing for Wireless Multimedia Sensor Networks. *Sensors* **2020**, *20*, 6217. [CrossRef]
17. Sun, F.; Xiao, D.; He, W.; Li, R. Adaptive Image Compressive Sensing Using Texture Contrast. *Int. J. Digit. Multimedia Broadcast.* **2017**, *2017*, 3902543. [CrossRef]
18. Gao, X.; Zhang, J.; Che, W.; Fan, X.; Zhao, D. Block-Based Compressive Sensing Coding of Natural Images by Local Structural Measurement Matrix. In Proceedings of the 2015 Data Compression Conference, Snowbird, UT, USA, 7–9 April 2015; pp. 133–142. [CrossRef]
19. Yu, Y.; Wang, B.; Zhang, L. Saliency-Based Compressive Sampling for Image Signals. *IEEE Signal Process. Lett.* **2010**, *17*, 973–976. [CrossRef]
20. Zhang, Z.; Bi, H.; Kong, X.; Li, N.; Lu, D. Adaptive compressed sensing of color images based on salient region detection. *Multimedia Tools Appl.* **2019**, *79*, 14777–14791. [CrossRef]
21. Feng, W.; Zhang, J.; Hu, C.; Wang, Y.; Xiang, Q.; Yan, H. A novel saliency detection method for wild animal monitoring images with WMSN. *J. Sens.* **2018**, *2018*, 3238140. [CrossRef]
22. Liu, W.; Liu, H.; Wang, Y.; Zheng, X.; Zhang, J. A novel extraction method for wildlife monitoring images with wireless multimedia sensor networks (WMSNs). *Appl. Sci.* **2019**, *9*, 2276. [CrossRef]
23. Liu, G.; Zheng, X. Fabric defect detection based on information entropy and frequency domain saliency. *Vis. Comput.* **2020**, *37*, 515–528. [CrossRef]
24. Li, R.; Duan, X.; Guo, X.; He, W.; Lv, Y. Adaptive Compressive Sensing of Images Using Spatial Entropy. *Comput. Intell. Neurosci.* **2017**, *2017*, 9059204. [CrossRef]
25. Xin, L.; Junguo, Z.; Chen, C.; Fantao, L. Adaptive sampling rate assignment for block compressed sensing of images using wavelet transform. *Open Cybern. Syst. J.* **2015**, *9*, 683–689. [CrossRef]
26. Zhang, J.; Xiang, Q.; Yin, Y.; Chen, C.; Luo, X. Adaptive compressed sensing for wireless image sensor networks. *Multimedia Tools Appl.* **2016**, *76*, 4227–4242. [CrossRef]
27. Zhao, H.H.; Rosin, P.L.; Lai, Y.-K.; Zheng, J.-H.; Wang, Y.-N. Adaptive gradient-based block compressive sensing with sparsity for noisy images. *Multimedia Tools Appl.* **2019**, *79*, 14825–14847. [CrossRef]
28. Zhao, H.H.; Rosin, P.L.; Lai, Y.K.; Zheng, J.H.; Wang, Y.N. Adaptive block compressive sensing for noisy images. In *International Symposium on Artificial Intelligence and Robotics*; Springer: Berlin/Heidelberg, Germany, 2018; pp. 389–399.
29. Canh, T.N.; Dinh, K.Q.; Jeon, B. Edge-preserving nonlocal weighting scheme for total variation based compressive sensing recovery. In Proceedings of the 2014 IEEE International Conference on Multimedia and Expo (ICME), Chengdu, China, 14–18 July 2014; pp. 1–5. [CrossRef]
30. Li, R.; Duan, X.; Lv, Y. Adaptive compressive sensing of images using error between blocks. *Int. J. Distrib. Sens. Netw.* **2018**, *14*, 1550147718781751. [CrossRef]
31. Monika, R.; Samiappan, D.; Kumar, R. Underwater image compression using energy based adaptive block compressive sensing for IoUT applications. *Vis. Comput.* **2020**, *37*, 1499–1515. [CrossRef]
32. Gambhir, D.; Rajpal, N. Edge and Fuzzy Transform Based Image Compression Algorithm: EdgeFuzzy. In *Artificial Intelligence and Computer Vision*; Studies in Computational Intelligence; Lu, H., Li, Y., Eds.; Springer: Cham, Switzerland, 2017; Volume 672. [CrossRef]
33. Wang, J.; Wang, W.; Chen, J. Adaptive Rate Block Compressive Sensing Based on Statistical Characteristics Estimation. *IEEE Trans. Image Process.* **2021**, *31*, 734–747. [CrossRef]

34. Ferrigno, L.; Marano, S.; Paciello, V.; Pietrosanto, A. Balancing computational and transmission power consumption in Wireless Image Sensor Networks. In Proceedings of the 2005 IEEE International Conference on Virtual Environments, Human-Computer Interfaces and Measurement Systems, Messina, Italy, 18–20 July 2005; pp. 61–66.
35. Kazemi, V.; Shahzadi, A.; Bizaki, H.K. Multifocus image fusion using adaptive block compressive sensing by combining spatial frequency. *Multimedia Tools Appl.* **2022**, *81*, 15153–15170. [CrossRef]
36. Liu, Y.; Wang, L.; Cheng, J.; Li, C.; Chen, X. Multi-focus image fusion: A Survey of the state of the art. *Inf. Fusion* **2020**, *64*, 71–91. [CrossRef]
37. Mathews, J.; Nair, M.S. Adaptive block truncation coding technique using edge-based quantization approach. *Comput. Electr. Eng.* **2015**, *43*, 169–179. [CrossRef]
38. Kodak Lossless True Color Image Suite. Available online: <http://www.r0k.us/graphics/Kodak> (accessed on 11 October 2022).
39. CCTV Image of Leopard and Deer Images from Sanjay Gandhi National Park—Bing Images—Search. Available online: <https://www.bing.com/images/feed> (accessed on 1 February 2023).
40. Available online: <https://github.com/ChaofWang/Awesome-Super-Resolution/blob/master/dataset.md> (accessed on 10 December 2022).
41. Wang, Z.; Bovik, A.C.; Sheikh, H.R.; Simoncelli, E.P. Image Quality Assessment: From Error Visibility to Structural Similarity. *IEEE Trans. Image Process.* **2004**, *13*, 600–612. [CrossRef]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

Article

Integration of the Wang & Mendel Algorithm into the Application of Fuzzy Expert Systems to Intelligent Clinical Decision Support Systems

Manuel Casal-Guisande ^{1,2,*}, Jorge Cerqueiro-Pequeño ^{1,2}, José-Benito Bouza-Rodríguez ^{1,2}
and Alberto Comesaña-Campos ^{1,2,*}

¹ Department of Design in Engineering, University of Vigo, 36208 Vigo, Spain; jcerquei@uvigo.es (J.C.-P.); jbouza@uvigo.es (J.-B.B.-R.)

² Design, Expert Systems and Artificial Intelligent Solutions Group (DESAINS), Galicia Sur Health Research Institute (IIS Galicia Sur), SERGAS-UVIGO, 36213 Vigo, Spain

* Correspondence: manuel.casal.guisande@uvigo.es (M.C.-G.); acomesana@uvigo.es (A.C.-C.)

Abstract: The use of intelligent systems in clinical diagnostics has evolved, integrating statistical learning and knowledge-based representation models. Two recent works propose the identification of risk factors for the diagnosis of obstructive sleep apnea (OSA). The first uses statistical learning to identify indicators associated with different levels of the apnea-hypopnea index (AHI). The second paper combines statistical and symbolic inference approaches to obtain risk indicators (*Statistical Risk* and *Symbolic Risk*) for a given AHI level. Based on this, in this paper we propose a new intelligent system that considers different AHI levels and generates risk pairs for each level. A learning-based model generates *Statistical Risks* based on objective patient data, while a cascade of fuzzy expert systems determines a *Symbolic Risk* using symptom data from patient interviews. The aggregation of risk pairs at each level involves a fuzzy expert system with automatically generated fuzzy rules using the Wang-Mendel algorithm. This aggregation produces an *Apnea Risk* indicator for each AHI level, allowing discrimination between OSA and non-OSA cases, along with appropriate recommendations. This approach improves variability, usefulness, and interpretability, increasing the reliability of the system. Initial tests on data from 4400 patients yielded AUC values of 0.74–0.88, demonstrating the potential benefits of the proposed intelligent system architecture.

Keywords: design; machine learning; expert systems; fuzzy logic; automatic rule generation; information fusion; intelligent system; decision-making; Wang–Mendel

MSC: 68T27; 68T30; 68T37

1. Introduction

Intelligent systems are now a reality, present in numerous and diverse environments, both domestic and commercial, and increasingly accepted and used by society [1–12]. This work is framed in the health sector, where intelligent systems have reached a significant level of development and are increasingly present, being regularly integrated into hospital computing environments, allowing the improvement and facilitation of clinical decision processes, with all the benefits of improving the quality of services provided and reducing the associated healthcare costs. From the use of statistical learning models, both in their supervised [13–15] and unsupervised variants [14–18], to the use of models based on knowledge representation through expert systems [19–24], intelligent systems have incorporated different inference mechanisms to increase their usefulness as diagnostic support tools [1,2,5,10,25]. In this sense, the authors of this article have presented several other works and applications of intelligent clinical decision support systems (ICDSS) [3–6,10], of which the last two proposals, which aim to help in the diagnosis of obstructive sleep

apnea (OSA) [1,2] and are taken as a starting point to elaborate the proposal presented in this work, deserve to be highlighted.

OSA is a major respiratory disorder affecting approximately one thousand million people worldwide [26], most of whom are undiagnosed. It has symptoms that also occur in the general population. It is characterized by the repeated total or partial collapse of the upper airway during sleep, which has a significant negative impact on those who suffer from it. Once the medical team is faced with a suspected case, the usual diagnostic procedure is to perform specific sleep tests, such as cardiorespiratory polygraphy [27–29] and polysomnography [30–35]. The availability of these tests is often limited, and they are usually expensive. Among the various measures obtained from these tests, the apnea-hypopnea index (AHI) [35], which is the ratio of the number of apnea and hypopnea events experienced by the patient during the night to the total number of hours of sleep [1], should be highlighted.

In this context, and after introducing the pathology and how it is diagnosed, we will briefly comment on the previously mentioned ICDSS used for the diagnosis of OSA. The first of these [1] proposes a system that, starting from a dataset related to the patient's health profile (anthropometric information, habits, comorbidities, and medication use), and through the concurrent use [1,2,4,5,8,10–12] of a series of machine learning classification algorithms, as well as a correcting block based on the sequential use of the adaptive neuro-based fuzzy inference system (ANFIS) and a specific heuristic algorithm, makes it possible to calculate a set of indicators associated with different AHI levels, which, after proper interpretation, makes it possible to determine whether a patient could suffer from OSA and to estimate its severity. The second of the aforementioned works [2] proposes a new approach based on two heterogeneous sets of information; on the one hand, the information related to the patient's health profile, already commented on in the first work; and on the other hand, the information related to the symptoms reported by the patients themselves using a specific OSA questionnaire. Unlike in the first work, here the information related to the patient's health profile is processed by a single machine learning classification algorithm associated with a single AHI level, while the second dataset—related to the symptomatology—is processed by a set of cascaded expert systems supported by the use of fuzzy inference engines. These two types of processing allow two different risk metrics to be obtained, which are later combined by means of a utility function to determine a new metric that allows the risk of a patient suffering from an OSA to be assessed.

It is clear that both of these systems can be very useful when a medical team is faced with a suspected OSA case, helping them to differentiate between patients who may suffer from the disease. However, beyond the clear benefits from a clinical perspective, it is necessary to point out the strengths and weaknesses of each of the proposals based on their ability to formalize and diversify knowledge and manage uncertainty. These issues are presented in detail in Table 1 for each of the systems.

After analyzing the data in Table 1, it can be seen that the first approach addresses a partial formalization of knowledge. In terms of diversification, both approaches have the same objective, although it could be considered that the first one provides more diverse information, pointing to different AHI threshold levels, thus facilitating the determination of the severity of the patients' condition. In the second approach, the formalization of knowledge is improved by incorporating a cascade of expert systems. Similarly, and in contrast to the first approach, the use of statistical and non-statistical approaches allows for a more complete management of uncertainty. However, in this second approach, given the coexistence of a pair of risks of different natures that try to represent the same phenomenon, obtained as outputs respectively from the machine learning algorithm and the cascade of expert systems, as a prior step to the generation of recommendations, it is necessary to carry out their union or aggregation, in this case by applying an analytical function that, in a certain way, increases the uncertainty present in the process. Thus, the aggregation of these risk terms is a difficult task in itself, requiring not only explicit knowledge of their nature and meaning but also a qualitative and quantitative assessment of their influence on

the subsequent aggregation. In order to optimize this aggregation, there are many different works that propose different aggregation models based on the interpretation of the terms, from simple weighted sums [36–39] to intelligent systems, plus a variety of aggregation operators from different origins [2–4,10,12].

Table 1. Strengths and weaknesses of previous approaches from the authors’ works and which were used as a grounding for the present proposal. Comparisons are made in terms of the capabilities to formalize and diversify knowledge as well as to deal with uncertainty.

	First Proposal [1]	Second Proposal [2]
Formalization and diversification of knowledge	Incorporates a correcting block based on the use of ANFIS and a specific heuristic algorithm, through which a partial formalization of knowledge is carried out. Considers several AHI threshold levels in the prediction, which facilitates the subsequent prioritization of the patients according to the severity of their condition by different types of health professionals.	Formalizes knowledge by using of a cascade of expert systems. Considers a single AHI level, which in some way reduces the subsequent performance of the system and allows only a single threshold to distinguish between patients who suffer from OSA and those who do not. This estimation might be used by the professionals to identify those patients who suffer from the condition.
Uncertainty management	Performs uncertainty management based on the use of statistical approaches, but without implicit processing of vagueness.	Uses statistical and non-statistical approaches, and therefore performs a more complete management of uncertainty and vagueness. The definition of the method for the union or aggregation of the system risks uses a utility function, which in some way involves increasing the uncertainty associated with the process.

In this paper, we propose an agile and novel solution to the problem of aggregating terms, in this particular case, risk metrics, when presented under the conditions of a fuzzy inference process. Such an inference model, starting from a knowledge base consisting of a set of fuzzy rules, that is, expressed by means of “IF . . . THEN . . . ” structures using linguistic qualifiers, is able to compute a prediction that can in fact be interpreted as the consequence of a logical combination of its antecedent terms. In this way, an aggregation is obtained that considerably reduces the imprecision of the expressions by which the aggregation could be represented as a prior step to the formalization of rules and at the same time reduces the uncertainty due to a lack of knowledge by being able to model the expression of the rules in the inference calculus structure itself.

Therefore, in this paper, a proof of concept is designed, developed, and carried out for an intelligent system aimed at predicting the severity of an OSA case represented by the AHI level, which, combining already discussed and published proposals, presents as its main novelty a risk aggregation model based on a set of fuzzy inference systems whose knowledge bases are autonomously computed from the study database, represented by the risk values obtained and their corresponding class labels.

Thus, in the first part, and taking into account previous developments by the authors [1,2], a prediction is made for two risk terms, named *Statistical Risk* and *Symbolic Risk*, associated with different inferential models, and obtained for a set of thresholds of AHI:

- For the processing of the data related to the patient’s health profile, that is, for the objective data, different machine learning classification algorithms are considered, working concurrently [1,2,4,5,8,10–12], each of which is associated with different AHI threshold levels as in the approach in the first work, and through which it is possible to obtain a set of risk indicators, called *Statistical Risks*.
- The data relating to the patient’s symptoms, that is, those of a subjective nature, are processed in a similar way to the approach in the second work, using a series of expert systems arranged in a cascade, the output of which determines an indicator, the *Symbolic Risk*, a common value for the patient, applicable to all AHI levels.

In a second part of this work, based on the results obtained in its first part, an inference system will be designed, defined, and developed to aggregate the risk terms. This combination/aggregation of the pair formed by each of the different statistical risks for each AHI level and the symbolic risk will be carried out precisely at each AHI level, using a set of fuzzy expert systems whose knowledge bases are not explicitly defined but are obtained automatically from the starting data set, represented by the calculated risks and using the corresponding classification labels. For this purpose, an algorithm is integrated into the program flow that allows to obtain fuzzy rules from a dataset, with the calculated risk pairs as antecedents and the class labels of each pair as consequents. This is the well-known Wang-Mendel algorithm [40]. Through this approach, it will be possible to automatically generate a set of rules that allow each specific case to be evaluated by combining these rules to obtain a prediction.

Motivation and Conceptual Approach

Regarding the first stage, the motivation of this work is to combine, in a single intelligent system, two previously developed approaches, already discussed and published by the authors [1,2]. To this end, the most differential and novel proposals will be selected, in this case the use of a set of threshold levels of the AHI with the generation of risk measures associated with predictor models based, on the one hand, on statistical learning (*Statistical Risk*) and, on the other, on the representation of knowledge (*Symbolic Risk*). Thus, it is proposed that the intelligent system calculate a pair of risks for each AHI threshold, which will improve prediction and diagnostic stratification. Both proposals, as mentioned above, are published separately, and their union is proposed in this paper, introducing an important novelty that affects the way of aggregating the generated risks, which is addressed in the second stage.

This second stage is based on the Wang-Mendel algorithm [40] and constitutes the main conceptual and theoretical basis for the novelty and technical relevance of the work presented above. The generation of risks reproduces inference models, both statistical learning and knowledge-based using fuzzy logic, which have already been presented and discussed. However, the approach of developing a fuzzy expert system whose knowledge base is automatically derived from a plausible representation of the same data that feeds the previous systems is a significant difference for which the authors have no further evidence in the literature in the field of study. The aggregation of the prediction results of statistical and symbolic classifiers is one of the difficulties inherent to the joint use of these models. Usually, the combined work of models based on statistical learning and symbolic inference is approached from the point of view of a model that combines both approaches within the definition of its architecture. Similarly, and in a more general perspective, the aggregation of results derived from models with inductive learning (statistical inference) and what could be considered analytical learning (deductive/symbolic inference) has usually been treated from the point of view of reviewing the theoretical domain of instances and its influence on the hypothesis search space [41]. In this line, generalized towards the definition of a dataset that can provide answers to symbolic and statistical models, the creation of hybrid intelligent systems is evolving, not only with the aim of improving the aggregation of their results but also with the objective of creating their own architectures and models. However, in the case of the work presented in this article, the hybridization is limited and does not intend to define a hybrid architecture but rather to improve the aggregation of results, albeit starting from the same set of initial instances. Theoretically, the proposal is based on the basic principles of statistical learning, which refer to the existence of a set of pre-labeled features. To this, it adds the ability to obtain a knowledge base from those. And this is where the real potential of the proposal lies. When we mention “features” we are always referring to data in its various forms and expressions, understood as quantitative measures of information, and always, within a knowledge base, particular manifestations of the set of relationships that those features themselves establish among themselves. Thus, while these relations, ordered and structured in the form of logical rules, constitute ontological and permanent knowledge, data are ephemeral and transitory and therefore, by definition,

cannot constitute a solid knowledge base. This is precisely the reason for the difficulty of creating knowledge bases and the main differentiation between statistical learning and models based on knowledge representation, in this case represented by fuzzy inference systems. While statistical inference uses only data in its learning to find hypotheses for a mathematical predictive model through a process of optimization, symbolic inference necessarily requires a logical knowledge base on which to support its reasoning process, be it with logical rules, probabilistic approaches, or fuzzy approaches. For this reason, the idea that one set of data can feed both inference processes is complex and, in many cases, chimerical. Is it possible to derive a logically coherent, ordered, and structured knowledge base from a set of scattered, statistically relevant, and conveniently labelled data?

In order to find an answer to this question, the concept of an expert system itself has to be reconsidered. In general, expert systems are highly dependent on their knowledge base, which makes them very difficult to use, especially given the absence of this base in the definition. This difficulty, especially notable in those systems considered to be of the first generation [42], is dealt with by the second generation of expert systems [43] that adopt strategies of identification and generation of heuristic rules through rudimentary learning processes. In line with the latter and with the emergence of what we could call third-generation expert systems [44], this article proposes to automate the creation of the knowledge base of a fuzzy expert system through the automated generation of rules from a set of labelled numerical data.

Thus, as we shall see, the answer to the question above is affirmative under a number of constraints, and this forms the basis of the Wang-Mendel algorithm.

1.1. Wang and Mendel's Method

Essentially, the method proposed by Wang and Mendel [40] allows the automatic generation of a set of rules on fuzzy sets from both numerical data, that is, input-output data pairs (for example: $(t_1, t_2; w)$), and from potential fuzzy rules proposed by experts. Both rules will be combined into a knowledge base from which inferences can be performed—as in the example before: $f(t_1, t_2) \rightarrow w$.

Figure 1 presents a diagram that aims to summarize the operation of the method proposed by Wang and Mendel, which consists of five stages as described below [40].

- Stage prior to the application of the method: before applying the Wang–Mendel algorithm, the dataset to be used is prepared by identifying the input and output variables. The range of each variable, that is, their maximum and minimum values, is also determined. The user must also select the type of membership function to use. The original proposal by Wang and Mendel envisaged the use of triangular membership functions. It is also necessary to define the value of N for each variable, which must be an integer greater than or equal to one. This value is used to determine the number of sections for each of the associated membership functions.
- Stage 1—Division of the input and output spaces into fuzzy regions: in this stage, for each of the input and output variables considered, the problem domain is divided into $2N + 1$ sections, in this case using triangular functions, as originally proposed in the paper of Wang and Mendel [40]. $2N + 1$ sections are added, since the goal is to perform a division of the domain of each of the variables in such a way that there is a central or intermediate section. Figure 2 shows an example for $N = 1$, with three segments of the membership function (L , M , and H) for the two input variables (t_1, t_2) and the output variable (w) of the previous example. As can be seen, and in line with the proposal of the original paper, there is an overlap of the triangles, so that if the top vertex of the central triangle has a maximum degree of membership, at the same point the vertices of the neighboring triangles have minimum degrees of membership.
- Stage 2—Generation of fuzzy rules from the input-output data pairs: once the sections of the membership function associated with each variable have been determined, the rules are to be generated. This is done by first determining the degrees of membership associated with each of the sections of the different functions for each of the different

lines of the initial dataset. For this purpose, the variables are fuzzified through their respective membership functions on the basis of the data available in the dataset. For example, in Figure 3, it can be seen that for a given observation of the variable t_1 , it has a degree of membership of 0.3 to Lt_1 and a degree of membership of 0.7 to Mt_1 . For variable t_2 , it is observed that it has a degree of membership of 0.45 to Mt_2 and a degree of membership of 0.55 to Ht_2 . In the case of w , it is observed that it has degrees of membership from 1 to Mw . After that, in each of the rows of the dataset, each variable is assigned to the section with the maximum degree of membership, determining a rule for each row. For the case in Figure 3, taking the maximum degrees of membership, the following would be obtained: $(t_1^{(1)}, t_2^{(1)}; w^{(1)}) \rightarrow \text{IF } t_1 \text{ is } Mt_1 \text{ and } t_2 \text{ is } Ht_2, \text{ THEN } w \text{ is } Mw$. As mentioned above, this process is carried out for each of the different rows of the dataset, with a rule being determined for each of these rows.

- Stage 3—Assignment of a degree to each of the rules to solve potential conflicts among the generated rules: the initial numerical set may include observations that, after applying Stage 2, generate rules that could be in conflict, that is, their antecedent part is the same while the consequent part is different. To deal with this problem, Wang and Mendel propose to associate each rule with a degree in order to select only those rules that have a maximum degree, thus dispensing with a large part of the rules generated in the previous stage. The degree value associated with each rule would be determined as the product of the degrees of membership of the observation that gave rise to the rule, as can be seen in Equation (1), which is particularized in Equation (2) for the case shown in Figure 3. In this case, after carrying out this process for all the rules and selecting those that maximize the coefficient value obtained, the resulting rule base would have the structure shown in Figure 4, with each of the boxes containing a section of the membership function of the consequent, if this combination were possible and present in the initial dataset.

$$D_{rule\ i} = \mu(x_1)_i \cdot \mu(x_2)_i \cdot \mu(y)_i \quad (1)$$

$$D_{rule\ 1} = \mu(t_1)_1 \cdot \mu(t_2)_1 \cdot \mu(w)_1 = 0.70 \cdot 0.55 \cdot 1 = 0.39 \quad (2)$$

- Stage 4—Building of the combined fuzzy knowledge base: once the knowledge base has been determined by the Wang–Mendel algorithm, whose structure is shown in Figure 4, it could be enriched by a set of fuzzy rules expressed by the expert team. These rules, which in Wang–Mendel’s definition are called linguistic rules or expert rules, would be incorporated into the knowledge base after being assigned a certain degree of importance by the expert team. In the case of conflict between the rule proposed by the expert and the rule generated by the algorithm in any of the boxes, Wang and Mendel advocate using the one with the maximum degree.
- Stage 5—Inference: once the knowledge base has been determined, it is possible to integrate it into an inference system, such as the Mamdani inference system [45–48], to draw conclusions with new input data.

This work is organized into five sections. This section introduces the background and context in which the work is developed. Next, the fundamentals of the method proposed by Wang and Mendel were presented. Section 2 deals with the conceptual design and implementation of the proposed system. Section 3 then presents a case study as a proof of concept, which aims to demonstrate how the system works. In Section 4, a discussion of the proposed architecture is presented. Finally, Section 5 addresses the main conclusions of this work.

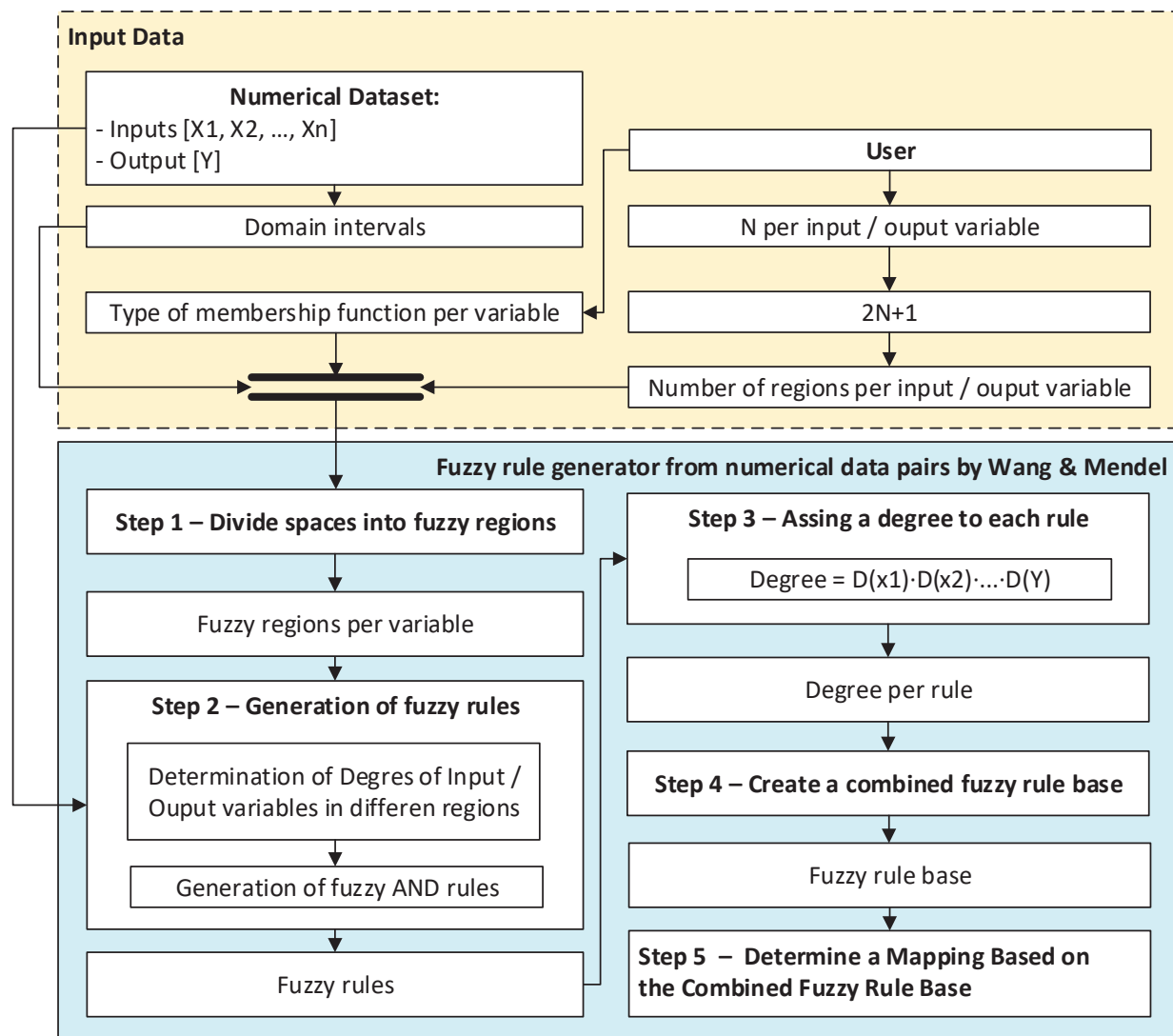


Figure 1. Diagram of the method for generating declarative rules on fuzzy sets as proposed by Wang and Mendel, showing each of the five stages over which the algorithm is deployed [40].

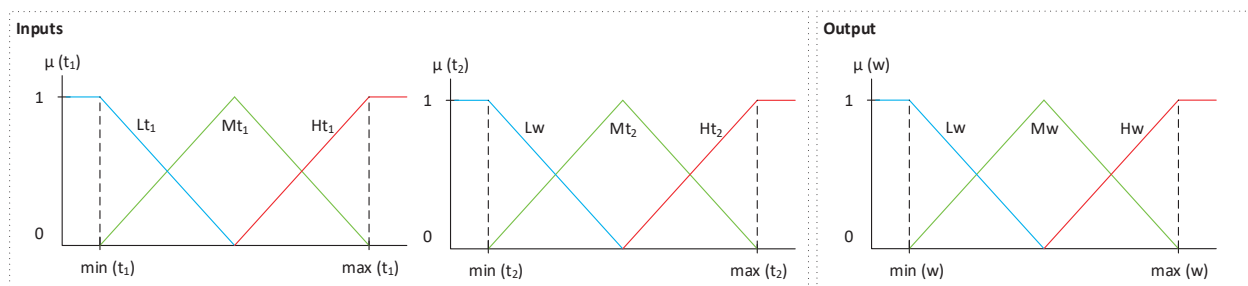


Figure 2. An example of the division of the initial domain for $N = 1$.

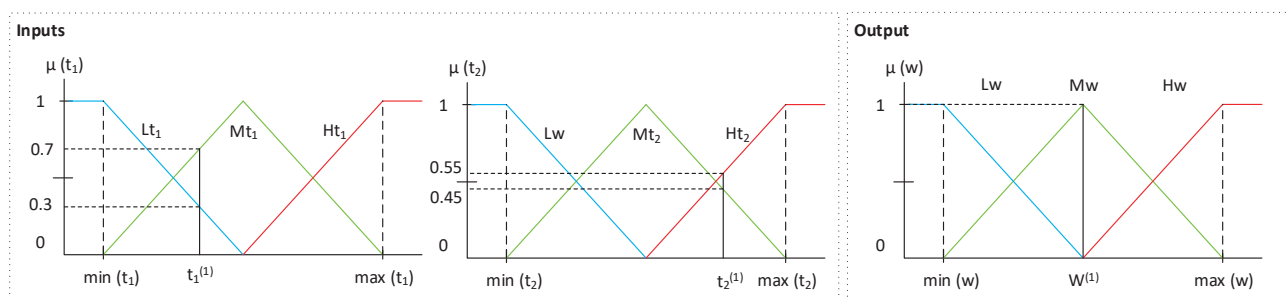


Figure 3. Example of the determination of degrees of membership based on the value obtained for the variables t_1 , t_2 and w in the initial dataset.

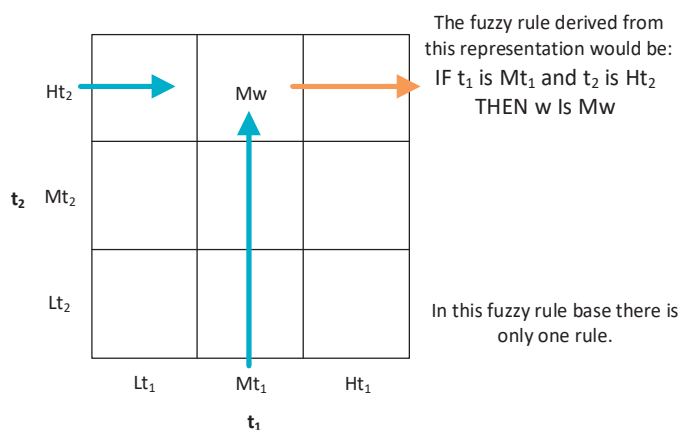


Figure 4. Distribution of the knowledge base generated by the Wang Mendel algorithm where the fuzzy rule generated in the example can be observed.

2. Materials and Methods

2.1. Definition of the System

2.1.1. Database Usage

This work is based on a database of 4583 patients, with information collected between the years 2013 and 2022 at the Sleep Respiratory Diseases Unit of the Pulmonary Department of the Álvaro Cunqueiro Hospital in Vigo (Galicia, Spain).

The information contained in the database related to patients suspected of suffering from OSA, can be divided into two groups. On the one hand, there is the information collected directly by expert pulmonologists, which is supposed to be accurate and unquestionable, and which refers to general patient data (gender, age, body-mass index and neck perimeter), their habits (tobacco and alcohol consumption), diagnosed pathologies conditions (hypertension, resistant hypertension, acute cerebrovascular accident (ACVA), ACVA less than a year ago, diabetes mellitus, ischemic heart disease, chronic obstructive pulmonary disease (COPD), home oxygen therapy, rhinitis, depression, atrial fibrillation and heart failure), and medication taken by the patient (benzodiazepines, antidepressants, neuroleptics, antihistamines, morphics and tranquilizers/hypnotics). As mentioned above, the information derived from the work of the expert pulmonologists is understood to be of low uncertainty and imprecision as it is validated and thoroughly reviewed. Otherwise, it would be an uncertainty factor that would need to be considered and estimated for control by the proposed intelligent system.

Figure 5 shows, in blue, a detailed description of the different aforementioned sub-groups and variables, as well as a description of the nature and type of the variable—whether it is a numerical or categorical variable. On the other hand, there is the information provided by the patient during a specific OSA interview (hours of sleep, minutes taken to fall asleep, prolonged intra-sleep awakenings, feeling of unrefreshing sleep, daytime tired-

ness, morning dullness, snorer, high intensity snorer, snore-related awakenings, unjustified multiple awakenings, nocturia, breathlessness awakenings, and reported apneas). This information is organized into different subgroups (sleep time subgroup, unrefreshing sleep subgroup, complicating sleep factors subgroup, snores subgroup), which are highlighted in orange in Figure 5. Additionally, the type of each variable nature, numerical or categorical, is indicated. Figure 5 shows all the initial information, both based on objective data, the use of which will be the focus of the machine learning models, and subjective data, from which the knowledge base of the expert systems that make up the symbolic part of the system will be elaborated.

General and anthropometric data	Data type	Pharmacological treatments	Data type
Gender	Categorical	Benzodiazepines	Categorical
Age	Numerical	Antidepressants	Categorical
Height	Numerical	Neuroleptics	Categorical
Body mass	Numerical	Antihistamines	Categorical
Body mass index (BMI)	Numerical	Morphic	Categorical
Neck circumference length (NCL)	Numerical	Relaxing/hypnotic drugs	Categorical
Habits of the patient	Data type		
Smoker	Categorical	Sleep time	Data type
Cigarettes per day	Numerical	Hours of sleep	Numerical
Years as a smoker	Numerical	Minutes until falling asleep	Numerical
Packs-per-year index	Numerical	Prolonged intra-sleep awakenings	Categorical
Drinking habits	Categorical	Unrefreshing sleep	Data type
Grams of alcohol	Numerical	Feeling of unrefreshing sleep	Categorical
		Daytime tiredness	Categorical
Illnesses	Data type	Morning dullness	Categorical
Hypertension	Categorical	Complicating sleep factors	Data type
Resistant hypertension	Categorical	Unjustified multiple awakenings	Categorical
ACVA	Categorical	Nocturia	Categorical
ACVA less than a year ago	Categorical	Breathless awakenings	Categorical
Diabetes	Categorical	Reported apneas	Categorical
Ischemic heart disease	Categorical	Snores	Data type
COPD	Categorical	Snorer	Categorical
Need for home oxygen-therapy	Categorical	High intensity snorer	Categorical
Rhinitis	Categorical	Snore related awakenings	Categorical
Depression	Categorical		
Atrial fibrillation	Categorical	Cardiorespiratory polygraphy	Data type
Heart failure	Categorical	Apnea-hypopnea index (AHI)	Numerical

Objective data
Subjective data
Cardiorespiratory polygraphy output

Figure 5. Structure of the database showing the different natures of the initial data, distinguishing those that are more objectifiable through objective measurements or responses from those that are more subject to interpretation by the pulmonary medical team and the patients themselves.

In addition to this information, the database also contains information related to the specific sleep tests that were performed on the patient (mainly cardiorespiratory polygraphs). In this sense, the apnea-hypopnea index (AHI) stands out, shown in green color in Figure 5.

From the initial dataset of 4583 patients, 183 lines were extracted and excluded from the training and validation process of the system, being reserved for a later test associated with the proof of concept of the proposed intelligent system. In this sense, it will be possible to carry out an independent analysis of the proposal presented, which will make it possible to highlight its relevance and applicability in a practical way.

2.1.2. Conceptual Design and Description of the System

The intelligent system has been designed with a sequential structure where, in the first step, the initial data are collected, then in the second step, the risks are determined, and finally, in the third step, they are aggregated to obtain the final prediction. Next, Figure 6 shows the flow diagram of the ICDSS proposed in this work, which will be described in detail below.

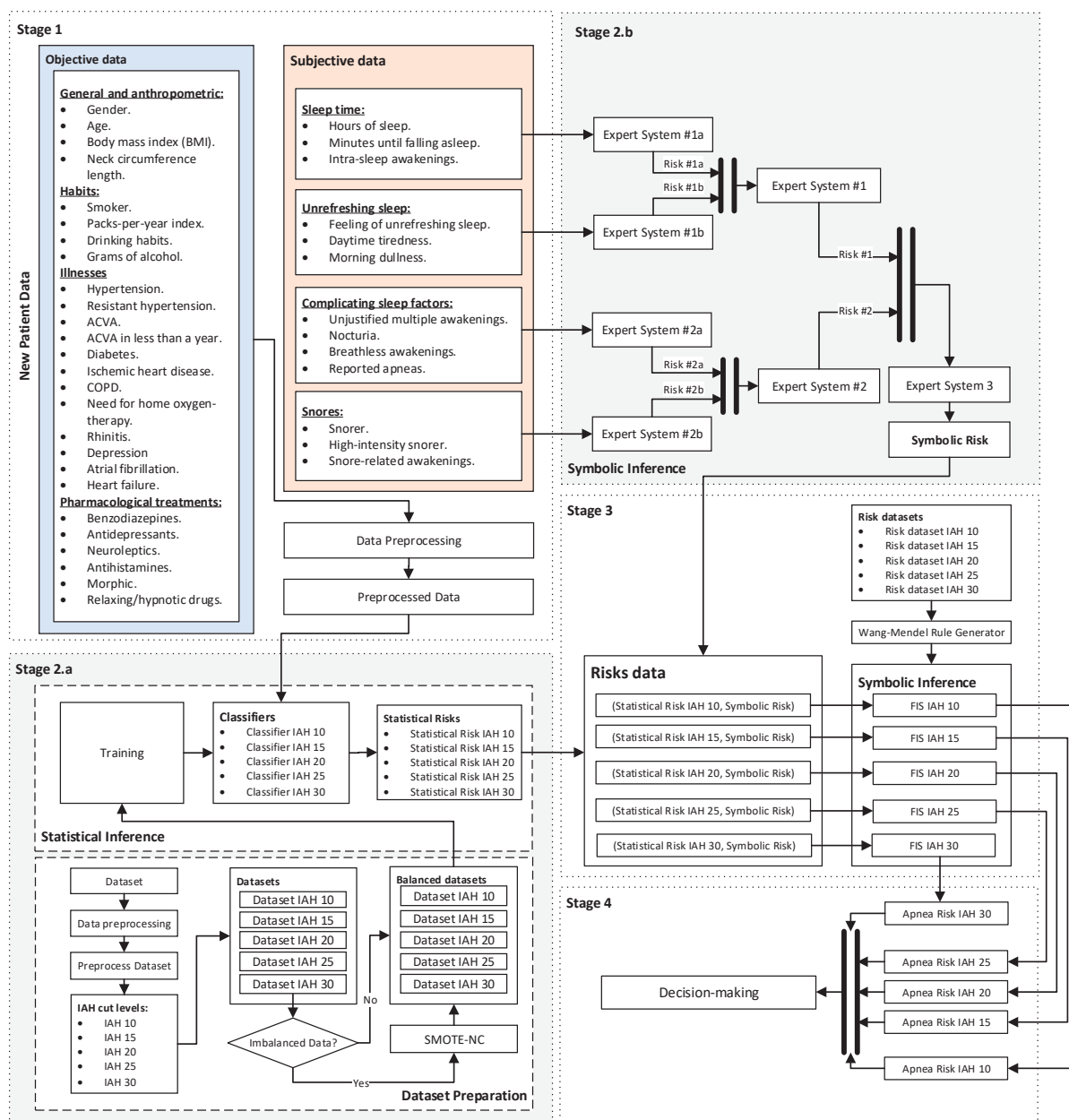


Figure 6. Flowchart of the Intelligent Clinical Decision Support System, showing how the information progresses through three stages. In Stage 1, the initial information is collected and pre-processed; in Stage 2.a, the information compiled by the expert pulmonologists (general information, habits, diseases suffered by the patient and medical drugs taken) is processed; in Stage 2.b the processing of the information related to the symptoms referred by the patient is carried out; in Stage 3, the automated rule generation approach proposed by Wang–Mendel is applied, defining a series of knowledge bases that allow to perform the combination of the risks obtained in the previous stages, making possible to determine a final risk value for each AHI level; and finally, in Stage 4, the risk evaluation and the corresponding decision making are carried out.

Stage 1: Collection of Patient Information

As shown in Figure 6, the first stage of the ICDSS focuses on the collection of the patient's starting information, both that information collected by the expert pulmonologists team and that reported by the patient using a specific OSA questionnaire, in which their symptoms are assessed in line with that mentioned in Section 2.1.1. Figure 5 may be visited for more information about the variables.

Stage 2: Determination of the Values of *Statistical Risk* and *Symbolic Risk* Values

Once the patient information has been collected and structured, it is processed using two sub-stages that run concurrently. These stages are based on those originally developed for the aforementioned papers by the authors [1,2], which constitute the basis for this proposal. Such a proposal addresses, on the one hand, the determination of a risk prediction associated with the interpretation of the patient's objective data by machine learning algorithms. On the other hand, a risk will also be calculated, this time associated with the generation of a knowledge base derived from the interpretation of the subjective data by the team of expert pulmonologists, based on their experience. This knowledge base will form part of the set of expert systems that will determine the aforementioned risk.

The first of these, sub-stage 2.a, focuses on the processing of the information collected by the expert pulmonologists, highlighted in blue in Figure 5, using a set of machine learning algorithms that work concurrently [1,2,4,5,8,10–12], through which it is possible to determine a set of scores, referred to in this work as *Statistical Risks*. For the definition and configuration of those algorithms, starting from the data presented in Section 2.1.1, a series of training datasets are built based on different AHI threshold levels (10, 15, 20, 25, and 30), so that in each of them two classes are defined, *OSA case* and the *non-OSA case*. If considered, and if there are medical reasons that justify it, additional threshold levels could be incorporated. All the development and technical details are described in the authors' previous works, in particular in the 2023 paper entitled "Design and Conceptual Proposal of an Intelligent Clinical Decision Support System for the Diagnosis of Suspicious Obstructive Sleep Apnea Patients from Health Profile" [1].

On the other hand, in the second of these stages, 2.b, the processing of the information corresponding to the symptomatology, which has a more subjective character and is highlighted in orange in Figure 5, is carried out using a series of expert systems working in cascade, based on Mamdani-type fuzzy inference engines [45–48]. Furthermore, in Figure 7, a detail of that cascade is shown.

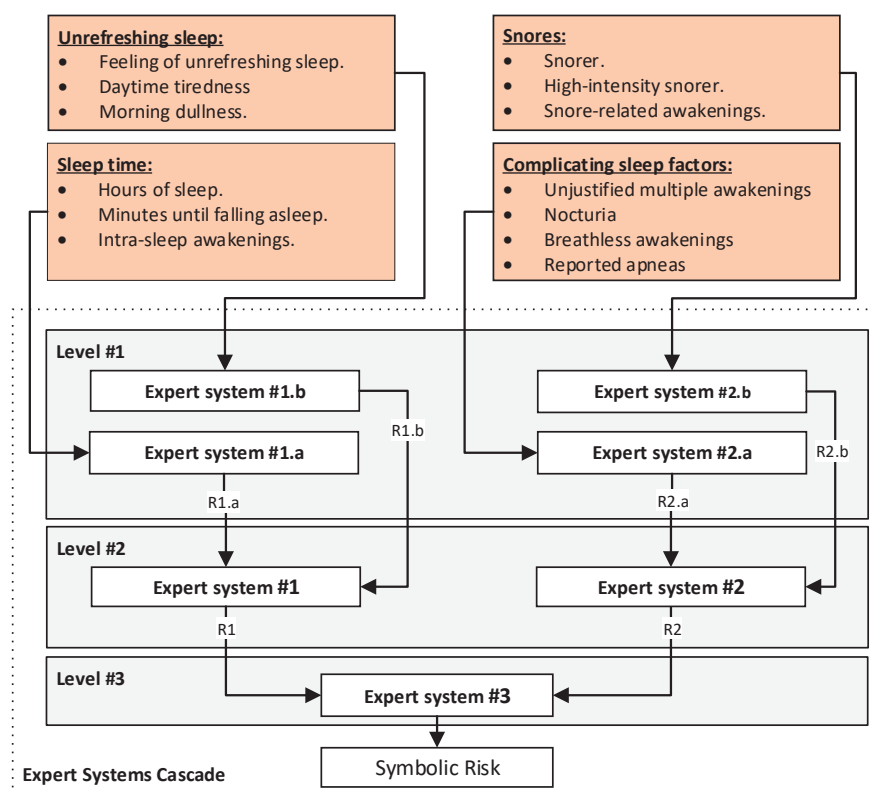


Figure 7. The cascade of expert systems in detail, showing the different levels the inferential reasoning process has been divided into.

As can be seen in Figure 7, at each level of the cascade, different risk indicators are obtained, each related to the risk of suffering from OSA. These indicators are combined and grouped at successive levels, finally determining at the output of the cascade a risk indicator that groups them, called *Symbolic Risk*. This risk is a common and general metric associated with the patient's own likelihood of developing OSA, and it is therefore not associated with any AHI threshold. As in the previous stage, the technical details of this proposal have already been established in a previous work by the authors [2].

Stage 3: Determination of the *Apnea Risk* Level for Each AHI Threshold Level

Stage 3 is key to the design of this system, as it contains the main contribution and novelty of the system. Once the indicators have been determined, both the *Statistical Risks* derived from statistical learning, and the result of the cascade of expert systems, that is, the *Symbolic Risk*, an objective aggregation must be considered in order to determine a single value related to the risk of a patient suffering from OSA.

Among the different aggregation models in this work, we chose to use a fuzzy expert system because of its capabilities to formalize and diversify knowledge, which would allow us to establish plausible reasoning about how to link previous risk indicators to explain a final state of risk. This would make it possible to gain explanatory power in the aggregation, which would undoubtedly help the medical expert team in interpreting the system's suggestions.

Therefore, once the different pairs (*Statistical Risk*, *Symbolic Risk*) have been determined for each AHI level, they are aggregated using a series of expert systems with Mamdani-type fuzzy inference engines [45–48]. By means of these systems, at each AHI level, it is possible to determine a risk metric that groups and represents them, the *Apnea Risk*.

However, the main difficulty in developing an expert system lies in identifying and creating its knowledge base. Without this knowledge base, the inference engine cannot work, and therefore the intelligent system would not work either. In this case, the knowledge base for each AHI level must consider how to aggregate the various risk indicators previously obtained to obtain an apnea global risk level. However, it is clear that the expert team has no explicit knowledge or experience in aggregating these terms, which are new concepts associated with their respective prediction models. Even considering that, obvious rules such as a joint ratio of high *Statistical* and *Symbolic Risk* values, obtaining a high *Apnea Risk*, could be questioned depending on the certainty of the initial data, the inferential process, and the lack of knowledge of the interpretation of the “high” fuzzy set. This lack of experience forces the medical team to assume that risk aggregation is an unexplained, statistical, and somewhat stochastic process. However, the chosen aggregation process is inherently explainable but difficult to implement in the absence of a knowledge base. How can this dilemma be tackled? This paper proposes the use of the Wang-Mendel algorithm to generate an explicit knowledge base of fuzzy rules from a data set. In other words, it addresses the traditional dilemma between data and knowledge present in any knowledge base that feeds an inferential symbolic engine.

Therefore, the elaboration of the knowledge base associated with each Mamdani fuzzy inference system discussed above will make use of a set of data by creating ordered input-output pairs and, with them, generate a set of fuzzy rules. Moreover, in this case, the algorithm cannot start from a set of existing linguistic rules. For this, it will be necessary to determine a set of datasets for each AHI level [40]. The explanatory variables in these datasets are the (*Statistical Risk*, *Symbolic Risk*) pairs for each AHI level, while the explained variable is a number, 0 or 1, depending on whether the patient has an AHI level below, equal to, or above the threshold level.

Stage 4: Generation of Alerts and Decision-Making

Given the data of a new patient, after determining and aggregating the different risk indicators, an *Apnea Risk* indicator is obtained for each AHI level, which belongs to a

continuous domain in the interval between zero and one, understood as the membership in the ‘suffering from apnea’ class.

In order to facilitate the interpretation of each *Apnea Risk* value, it has been decided to establish a risk threshold value for each AHI level based on a graphical optimization process similar to that used in the work by Casal-Guisande et al. [1], based on determining the threshold value at which the Matthews correlation coefficient [49–51] is maximized.

The medical team will select the AHI level they feel is most appropriate to consider that a patient may be suffering from OSA, and in light of that previously mentioned threshold, the system will generate the appropriate alerts and facilitate the decision-making processes.

2.2. Implementation of the System

In order to implement the ICDSS proposed in Section 2.1, which addresses everything from the collection of information related to the patient to the generation of alerts and decision-making, the process of building a software artifact is described below, which has been developed taking into account the recommendations and guidelines proposed by Hevner et al. [52,53], thus guaranteeing, if necessary, that it can be integrated into hospital information systems.

For the development and implementation of the software artifact, MATLAB® (R2022b, 326 MathWorks®, Natick, MA, USA) was used, as well as Python (version 3.9.12), together with a series of packages explained in Table 2. The software artifact is accompanied by a graphical user interface to facilitate interaction with it (see Figure 8).

The screenshot shows the 'Apnea Diagnosis - Wang Mendel Version App' window. It features a 'See WM Knowledge Bases' button at the top left. The main interface is organized into three primary sections, each highlighted with a colored border: Region (1) in red, Region (2) in blue, and Region (3) in purple.

- Region (1) Initial data collection (1):** This section is divided into 'Objective data' and 'Subjective data'.
 - Objective data:** Includes 'General and anthropometric data' (Gender, Age, Mass, Height, BMI, Neck circumference), 'Habits' (Smoker, Cigarettes per day, Years as a smoker, Packs-per-year index, Drinking habits, Grams of alcohol), 'Pharmacological treatments' (Antidepressants, Antihistamines, Benzodiazepines, Morphic, Neuroleptics, Relaxing/hypnotic drugs), and 'Illnesses' (Hypertension, COPD, Resistant hypertension, Need for home oxygen-therapy, ACVA, Rhinitis, ACVA in less than a year, Depression, Diabetes, Atrial fibrillation, Ischemic heart disease, Heart failure).
 - Subjective data:** Includes 'Sleep time - R1.a' (Hours of sleep, Minutes until falling asleep, Prolonged intra-sleep awakenings), 'Unrefreshing sleep - R1.b' (Feeling of unrefreshing sleep, Daytime tiredness, Morning dullness), 'Complicating sleep factors - R2.a' (Nocturia, Breathless awakenings, Reported apneas), and 'Snoring - R2.b' (Snorer, High intensity snorer, Snore related awakenings).
- Region (2) Data Processing (2):** This section includes 'Statistical risks [0,100]' (R.10 to R.30), 'Symbolic risk' (LEVEL 1 - [1, 9], LEVEL 2 - [1, 9], LEVEL 3 - [0, 100]), and 'Risk aggregation [0,1]' (AHI 10 to AHI 30). Each section has a 'Calculate' button.
- Region (3) Alert generation & Decision-making (3):** This section includes 'Generate recommendation' (Select AHI threshold, Calculate (4), Clean) and 'AHI level estimation' (10, 15, 20, 25, 30). A 'System State' box is also present.

Figure 8. Screenshot of the application. Region (1) refers to the collection of the starting information. Region (2) refers to the processing of data. Region (3) refers to the generation of alerts and decision-making.

Table 2. List of the software packages used for the implementation of the software artefact.

MATLAB	
Toolbox	Comments
App Designer [54]	Facilitates the development of the user graphical interface for the artifact.
Classification Learner [55]	Allows to perform the training and massive machine learning classification algorithms test.
Fuzzy Logic Toolbox [56]	Makes possible to implement the fuzzy logic-based inferential engines.
Python	
Package	Comments
Imbalanced-learn library [57]	Provides different tools for addressing classification problems when unbalanced datasets are available. In this case, the SMOTE-NC algorithm is used.

In Figure 8, three regions stand out. Region (1), highlighted in red, refers to the stage of collecting initial information, both objective and subjective. Region (2), highlighted in blue, includes the processing of the data, taking into account the previously commended stages 2 and 3. Finally, region (3), highlighted in purple, refers to stage 4 and is related to the display of alerts and the generation of recommendations.

2.2.1. Data Collection

First of all, as mentioned above, the information from the patient to be studied must be entered into the application using the different fields highlighted in the red box in Figure 8. It is recommended that the data be verified once it is entered into the application in order to correct any errors or omissions that may lead to an increase in inaccuracy.

2.2.2. Data Processing

Once the data has been entered into the application, it is processed by the ICDSS. To do this, and in line with what has already been commented on, there is a region in the graphical interface, highlighted in blue in Figure 8, which consists of three panels.

The first two panels display the results obtained after applying a series of machine learning algorithms, as well as a cascaded set of expert systems through which it is possible to obtain a series of risk indicator values (*Statistical Risks* and *Symbolic Risk*, respectively). These risk indicators are later aggregated using a series of Mamdani-type fuzzy inference systems [45–47], whose knowledge bases are determined using an automatic rule generation approach, and finally make it possible to determine the *Apnea Risk* value for each AHI level, shown in the third panel of the blue region in Figure 8.

However, prior to processing the data of a new patient, it is necessary to detail the implementation of the different calculation engines. To do this, data from 4400 patients was extracted from the initial dataset commented on in Section 2.1.1.

Machine Learning Algorithms

The generation of statistical risks associated with using machine learning classifier algorithms on the initial dataset is briefly described below. A more detailed explanation can be found in the authors' 2023 papers [1,2].

To define the machine learning classification algorithms, the most objective information is used, collected by expert pulmonologists and highlighted in blue in Figure 5. Most of the variables are of nominal or categorical type [58,59], so they are encoded using dummy encoding [1]. The remaining variables, those corresponding to numerical data (BMI, age, etc.), are rescaled between zero and one using the MIN-MAX normalization method [1]. This is because, in all cases, it is possible to define the minimum and maximum values between which each of the variables will move, based on medical criteria. Furthermore, considering different AHI thresholds (10, 15, 20, 25, and 30), it is possible to analyze the results presented by each of the patients in the training dataset, generating a set of *OSA case* or *non-OSA case* labels associated with each patient and each AHI level. In this way, a

set of labeled datasets is created for the different AHI thresholds, from which the different machine learning classification algorithms are trained.

It is important to note that in healthcare settings, it is common to have unbalanced datasets, that is, significant differences in the number of patients in the different classes. In this case, this phenomenon is also observed (see Table 3). To solve this problem, a common and widely used practice in healthcare [1,5,60] is the use of approaches and strategies for controlled data augmentation. Given the heterogeneous nature of the data considered, the use of the Synthetic Minority Over-Sampling Technique for Nominal Continuous (SMOTE-NC) is chosen, a variant of SMOTE [60,61] with the ability to handle both numeric and categorical data. A number of neighbors $k = 5$ is defined, and data are added until a total of 4000 patients are available in each of the classes for the different AHI threshold levels, as can be seen in Table 3.

Table 3. Summary of the distribution of classes for the different AHI threshold levels.

Threshold	Before SMOTE-NC			After SMOTE-NC		
	AHI < Threshold	AHI $\geq$ Threshold	Total	AHI < Threshold	AHI $\geq$ Threshold	Total
10	1227	3173	4400	4000	4000	8000
15	1707	2693	4400	4000	4000	8000
20	1726	2274	4400	4000	4000	8000
25	2072	1928	4400	4000	4000	8000
30	2365	1635	4400	4000	4000	8000

Once the different training sets have been defined, the classification machine learning algorithms are trained using a 5-fold cross-validation strategy to achieve optimal results in terms of both hyper-parameter optimization and the generalization capacity of the chosen learning model. Considering the results obtained in the previous work of the authors, Casal-Guisande et al. [1], the use of bagged trees is chosen.

Figure 9 shows a summary of the different ROC curves obtained for the different AHI levels using the bagged trees algorithm.

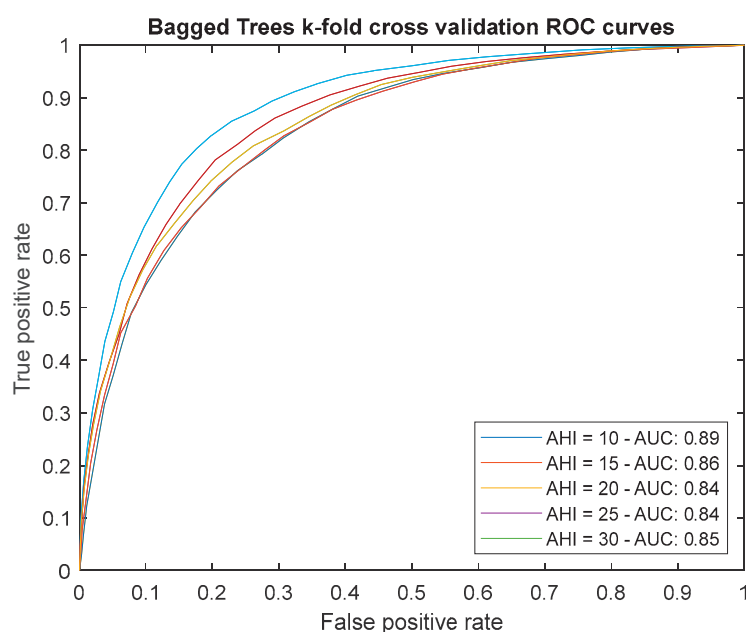


Figure 9. Plot of the ROC curves obtained for each of the AHI thresholds, taking into account the bagged trees algorithm and a 5-fold cross-validation. It can be seen that the area under the curve is greater than 0.8 in each case, with value 1 corresponding to a perfect classifier.

Cascaded Expert Systems

As in the previous section, the generation of the *Symbolic Risk* from a cascade model of expert systems is briefly described below. A more detailed explanation can be found in the authors' 2023 paper [2].

Concurrently [1,2,4,5,8,10–12] to the processing of the objective data by the machine learning classification algorithms for each level of AHI, the processing of the subjective data related to the symptoms reported by the patients is carried out by a series of expert systems arranged in a cascade. The output of this cascade is a risk metric value that will move between 0 and 100, called *Symbolic Risk*, related to the hazard of suffering from OSA in a general way, this time without this risk measure being associated with any AHI threshold value.

To do this, each expert system in the cascade uses a Mamdani-type fuzzy inference engine [23–25], similar to those used in the work by Casal-Guisande et al. [2] and others [4,5,7,8,11,12]. In line with what has already been commented on, the cascade of expert systems, shown in Figure 7, is distributed over three levels, as detailed in Table 4.

Table 4. Explanation about the sections of the cascade of expert systems considering the antecedents and consequents of each level [2].

Cascade Levels	Observations
Level 1	At the first level, four expert systems are used to process the information related to the symptoms reported by the patient (for more information, see Figures 6 and 7). The first of the expert systems is in charge of processing the set of information related to <i>sleep time</i> and determines the risk indicator <i>R1.a</i> at its output. The second of the expert systems focuses on processing the group of information related to <i>unrefreshing sleep</i> , determining the risk indicator <i>R1.b</i> at its output. The third of the expert systems focuses on the group of information related to <i>complicating sleep factors</i> , determining the risk indicator <i>R2.a</i> at its output. Finally, the fourth of the expert systems focuses on the <i>snORES</i> information group, determining the risk indicator <i>R2.b</i> at its output. Each one of the determined risk indicators is related to the group of information used for its determination and represents respectively the hazard level associated with suffering from an OSA case in relation to each group of data.
Level 2	Once the risk indicators have been determined at the first level of the cascade of expert systems, they are aggregated into groups of two (<i>R1.a</i> and <i>R1.b</i> , as well as <i>R2.a</i> and <i>R2.b</i>) using two new expert systems working concurrently [1,2,4,5,8,10–12] as can be seen in Figures 6 and 7, and at their output, after the defuzzification process, determine two new indicators, the risks <i>R1</i> and <i>R2</i> . These new risk indicators represent, respectively, the danger of suffering from OSA in relation to the risk indicators of the previous level, through which it was possible to determine each of the indicator.
Level 3	Finally, at the last level of the cascade, a final expert system aggregates the risk indicators obtained at the second level (<i>R1</i> and <i>R2</i>), determining at its output an indicator, called <i>Symbolic Risk</i> , related to the conjoint hazard that a patient has of suffering from an OSA case, after contemplating all the symptoms reported by the patient.

The use of the cascade, in addition to allowing the aggregation of the initial information, the reduction of the dimensionality of the problem, and the formalization of the knowledge, has associated advantages, such as greater simplicity in the process of elaboration of the rules since the number of antecedents to be considered in each inference system is less than if they were all considered at the same time, and therefore the process of elaborating and determining them is more precise. From a practical point of view, for the elaboration of the membership functions, the recommendation of Ross [48] was followed, opting for the use of normal, convex, and symmetric membership functions. Triangular and trapezoidal functions were used for the antecedents, and trapezoidal functions for the consequents. These choices are related to the nature of the data. Triangular functions are chosen when there is only one point at which the degree of membership is maximized; meanwhile, in the case of trapezoidal functions, there is a section—a range of values—in which the degree of membership is maximized.

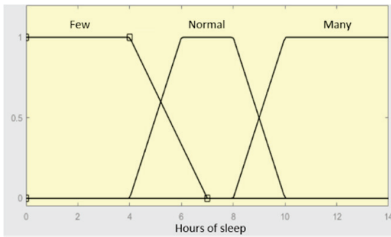
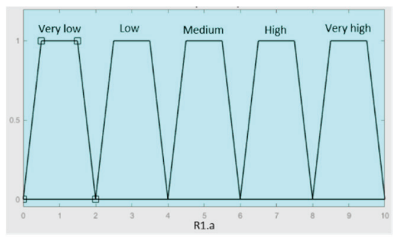
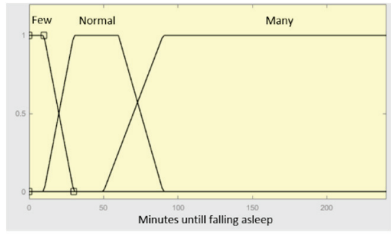
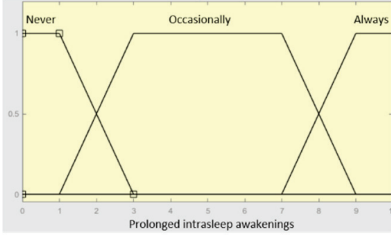
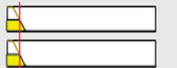
In a general way, Table 5 presents a summary of the configuration of the expert systems used in the cascade.

Table 5. General configuration of the expert systems used in the cascade [2].

Fuzzy Structure	Mamdani-Type
Defuzzification method	Centroid [48]
Implication method	Min
Aggregation method	Max

In order to explain in more detail, the configuration of one of the expert systems, Table 6 gives a detailed description of the expert system in charge of processing the data related to *sleep time*, from which the risk value *R1.a* is obtained. The remaining cascade expert systems are similar to this one, in line with what was commented in the work by Casal-Guisande et al. [2].

Table 6. Configuration of the expert system in charge of processing the *sleep time* data group [2].

Inference System Associated to the Sleep Time Data Group			
Input Data	Range	Output Risk	Range
Hours of sleep	0–14 h	R1.a	0–10
			
Minutes until falling asleep	0–240 min	Initial configuration Fuzzy structure: Mamdani-type. Membership function type: trapezoidal. Defuzzification method: centroid [48]. Implication method: MIN. Aggregation method: MAX. Number of fuzzy rules: 46	
			
Prolonged intra-sleep awakenings	0–10	Subset as an example of the 46 fuzzy rules 1. IF (<i>Hours_of_sleep</i> is <i>Few</i>) AND (<i>Minutes_until_falling_sleep</i> is <i>Few</i>) AND (<i>Prolonged_intra-sleep_awakenings</i> is <i>Never</i>) THEN (<i>R1.a</i> is <i>Low</i>). 2. IF (<i>Hours_of_Sleep</i> is <i>Few</i>) AND (<i>Minutes_until_falling_sleep</i> is <i>Few</i>) AND (<i>Prolonged_intra-sleep_awakenings</i> is <i>Never</i>) THEN (<i>R1.a</i> is <i>Medium</i>).	
			
Graphical example of fuzzy rules 1 and 2			
Hours of sleep = 3.52 1  2  Minutes until falling asleep = 20.5   Prolonged intrasleep awakenings = 1.13   Sleep periods 			

Risk Aggregation by Means of Fuzzy Expert Systems with an Automatically Generated Knowledge Base

After processing the initial data, both those having an objective nature and those having a subjective nature, a set of risk indicators is obtained, grouped into pairs (*Statistical Risk*, *Symbolic Risk*) for each AHI level.

At each AHI level, these indicators are processed using a Mamdani-type inference system [45–48], whose knowledge base is determined using an automatic rule generation approach, more specifically that proposed by Wang and Mendel [40]. At the output of each inference system, after aggregating the *Statistical Risk* and *Symbolic Risk* indicators, a final risk indicator associated with the hazard of suffering from OSA is obtained, called *Apnea Risk*.

Determination of Knowledge Bases Using Wang and Mendel Algorithm

As already mentioned, in order to determine the knowledge base at each AHI level, the algorithm proposed by Wang and Mendel [40] for the automatic generation of declarative rules on fuzzy sets will be used, the operation of which has already been detailed in Section 1.1. However, before doing so, it is necessary to comment on the structure of the datasets to be used.

It is appropriate to argue that the starting data that will feed the Wang and Mendel algorithm will be derived from those that make up the target dataset used in statistical learning algorithms. The latter is used because it is the one that should be considered more accurate and therefore with less uncertainty, as well as already being suitably labeled. The use of the augmented dataset is not considered because the algorithm does not have the appropriate adjustment and generalization capabilities, nor is the subjective data integrated into it because it is already part of a previous knowledge base that feeds the different levels of the cascade.

Therefore, starting from the training dataset before data augmentation, for each AHI level (10, 15, 20, 25, and 30), and after the processing by the machine learning algorithms as well as by the expert systems cascade, there will be a dataset with the values of *Statistical Risk* and *Symbolic Risk* and a label (0 or 1) for each patient and AHI threshold level. The label basically aims to indicate whether the patient has an AHI value greater than or equal to (a value of one) or less than (a value of zero) the threshold. Since the training dataset had 4400 patients, there will be a total of 4400 rows of data in each of the datasets for each AHI level.

Once the initial dataset has been defined, the automatic generation of rules at each of the AHI levels is carried out. To do this, we will follow the step-by-step structure discussed in the previous section for implementing the algorithm proposed by Wang and Mendel.

Stage prior to the application of the method:

Since the datasets to be used have an identical structure, the same strategy is followed in the different cases, dividing the initial spaces of each of the different variables using triangular functions, both for the antecedents (*Statistical Risk* and *Symbolic Risk*) and for the consequents (the label that represents the global risk of suffering from apnea). In addition to this, the value of N must be chosen, which, as already mentioned, is related to the number of sections that the membership function of each variable will have. In this case, it is decided to use $N = 4$ for the antecedents, giving a total of 9 sections. The choice of $N = 4$ for the case of antecedents is not arbitrary but refers to Miller's original work [62], which concluded that people can generally process information about seven events simultaneously, with a variation of plus-or-minus two in their number. This, applied to the fuzzification of a variable, suggests that membership functions with very few sections, or with a number of sections greater than nine, would be very difficult for a human to interpret and would therefore represent information that is difficult to express in a repetitive and common way using the usual language qualifiers. On the other hand, in the case of the consequent (the label), it is decided to use $N = 1$, giving three possible sections of the membership function. This is because the label represents only discrete values of zero or one, which belong only to the extreme sections, so adding more sections would not give any advantage.

Stage 1—Division of the input and output spaces into fuzzy regions:

Given that there are two antecedents, each one with nine sections in the membership functions, obtained by applying the calculation of the number of the sections as $2N + 1$, a grid of 81 rules is obtained, in line with what was commented in Section 1.1 and shown in Table 7 for an example case of AHI = 15. However, it should be noted that it will not always be possible to fill in the grid with the 81 rules, since it could happen that not all the possible combinations appear in the datasets used, either because these are not large enough or because cases that are not possible are represented. Once this is complete, the Wang–Mendel algorithm [40] is used.

Table 7. Example of grid of rules for AHI = 15. The variables *mfsy* represent the sections of the membership function of the antecedent associated with the *Symbolic Risk*; The variables *mfst* represent the sections of the membership function of the antecedent associated with the *Statistical Risk*; The variables *mf* represent the sections of the membership function of the consequent associated with the labels of the class ‘OSA case’ or ‘non OSA case’.

Statistical Risk	<i>mfst9</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>
	<i>mfst8</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>
	<i>mfst7</i>	-	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>
	<i>mfst6</i>	AR: <i>mf3</i>	-	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	-
	<i>mfst5</i>	-	-	-	AR: <i>mf1</i>	AR: <i>mf3</i>	AR: <i>mf3</i>	AR: <i>mf1</i>	AR: <i>mf3</i>	AR: <i>mf3</i>
	<i>mfst4</i>	-	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	-
	<i>mfst3</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>
	<i>mfst2</i>	-	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>
	<i>mfst1</i>	-	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>	AR: <i>mf1</i>
Apnea Risk (AR)	<i>mfsy1</i>	<i>mfsy2</i>	<i>mfsy3</i>	<i>mfsy4</i>	<i>mfsy5</i>	<i>mfsy6</i>	<i>mfsy7</i>	<i>mfsy8</i>	<i>mfsy9</i>	
Symbolic Risk										

Stage 2—Generation of fuzzy rules:

Once the set of antecedents and consequents has been established, the rules are automatically generated according to the values found in the initial data. It should be noted that only those rules are taken whose data reflect any membership in the previously created fuzzy sets of antecedents and consequents.

It should also be noted that in this case, the logical union of the antecedents in the rules is considered AND-logic since it is assumed that both antecedents are necessary to activate a rule and to determine the consequent.

Stage 3—Assignment of a degree to each of the rules:

Once the rules have been created, they are ordered according to their degree value, discarding those rules with lower degrees that share antecedents. The degree value shall be decided by using the product of the corresponding degree of membership of the values of each data line to the fuzzy sets represented in the corresponding rules, as reflected in Equations (1) and (2). If two rules share antecedents, the one with the higher degree is taken.

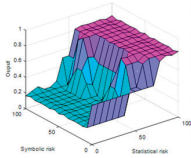
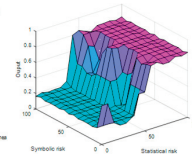
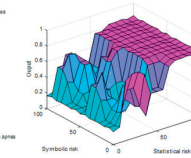
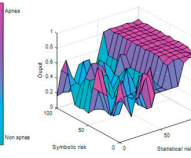
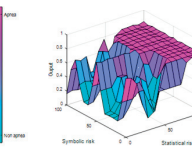
Stage 4—Building of the combined fuzzy knowledge base:

The Wang–Mendel algorithm makes it possible to combine fuzzy rules generated from a set of numerical data with rules defined by a team of experts. These latter, considered linguistic or expert rules, reflect the explicit reasoning in the domain of application of the algorithm. In this case, as mentioned above, the team of medical experts does not have the necessary and sufficient information to create a set of rules, so only the rules generated from the data are used.

Table 8 below gives a summary of the number of rules determined in each of the cases, that is, for each AHI level, as well as the generated surface through which the inputs

(*Statistical Risk*, *Symbolic Risk*) are mapped to the outputs (membership to the ‘suffering from apnea’ set, in the context of this work: *Apnea Risk*). In line with the comments above, it can be seen that the number of rules obtained is less than 81. This is due to the fact that not all the combinations of the antecedents are represented in the datasets used.

Table 8. Summary of the number of rules and the surfaces generated for each AHI level.

	AHI 10	AHI 15	AHI 20	AHI 25	AHI 30
Number of rules	68	71	72	72	71
Generated surface					

Stage 5—Inference:

Finally, we proceed to the inferential calculation of the aggregation results at each AHI level, considering a Mamdani-type fuzzy inference system and the different knowledge bases conveniently generated from each of the sets of fuzzy rules reflected in Table 8.

The knowledge bases, one for each AHI level, created above are based, as already argued, on the set of objective data, mainly represented by the *Statistical Risks* (variable antecedent) and Labels (consequent), which formed the starting dataset of the statistical learning models, including the influence of the *Symbolic Risk* (constant antecedent). The Wang-Mendel algorithm makes it possible to create a knowledge base based not on the traditional acquisition of this knowledge but on ephemeral data, which is undoubtedly a unique innovation. However, this ephemeral character of the data is transferred to the knowledge base, which could be updated according to the volatility of the initial data, sharing this transitory and harmful characteristic with a model based on the structured and permanent representation of knowledge. The medical expert team should review the generated knowledge base and assess whether the knowledge expressed in it should be consolidated so that it becomes permanent and is not subject to the volatility of the numerical input data. In addition, they will be given an explanation about the inference process itself and the logical reasoning implicit in it, and they will be able not only to evaluate the results of the aggregation but also to explain, transmit, and teach them, thus fulfilling the foundations of formalization and diversification that underlie the definition of any expert system.

Proof Test Results

After defining the knowledge bases that establish how to combine the *Statistical Risk* and the *Symbolic Risk*, it is proceeded to analyze the generalization capabilities of the system on an independent test dataset with 183 patients, different from that used in the construction of the ICDSS.

Figure 10 below shows the ROC curves obtained on the test dataset for each AHI level, with AUC values ranging from 0.74 to 0.88.

Determination of a Threshold Level for Each AHI Level

Once the results on the test dataset have been determined, in order to interpret the results obtained at the output of each of the inference systems that aggregate *Statistical Risk* and *Symbolic Risk*, and to understand them as binary classifiers [63,64], it is necessary to establish a cut-off level that allows discrimination between the ‘OSA case’ and ‘non OSA case’ classes. To do this, starting from the patients in the test dataset, an optimization process is carried out that aims to determine the cut-off value that, in each case, allows maximizing the Matthews correlation coefficient (*Mcc*) value (see Equation (3)) [49–51]. The acronyms

in the equation are as follows: TN = True Negatives, FN = False Negatives, TP = True Positives and FP = False Positives.

$$Mcc = \frac{TN \cdot TP - FN \cdot FP}{\sqrt{(TP + FP) \cdot (TP + FN) \cdot (TN + FP) \cdot (TN + FN)}} \quad (3)$$

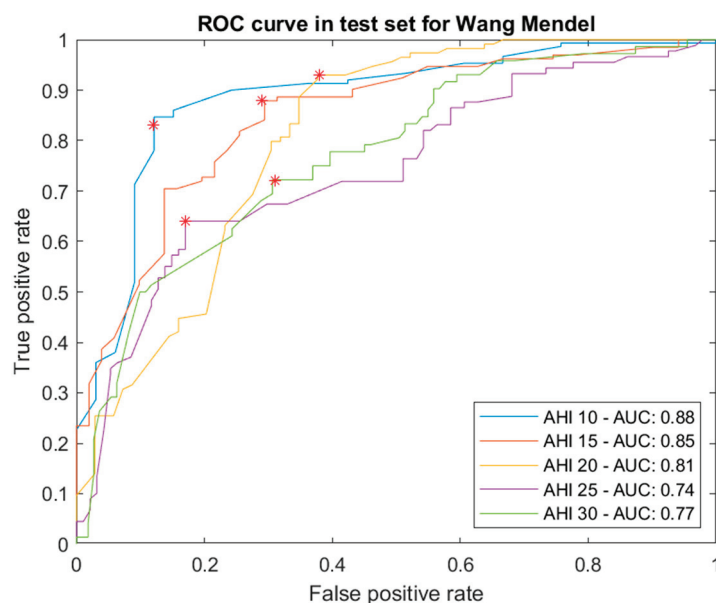


Figure 10. ROC curves analysed in the test set and derived from the classifiers of each level of the AHI index after risk aggregation. The red asterisk indicates the point that optimises the classification values obtained after the process of optimising the Matthews coefficient for each case. In all cases, reasonably good curves with AUC values above 0.74 are observed.

Figure 11 below shows the graphs of the Mcc values for the different thresholds in each AHI level. At the cut points in Figure 11, Mcc values of 0.60, 0.58, 0.60, 0.48 and 0.41 are obtained for AHI levels of 10, 15, 20, 25 and 30 respectively.

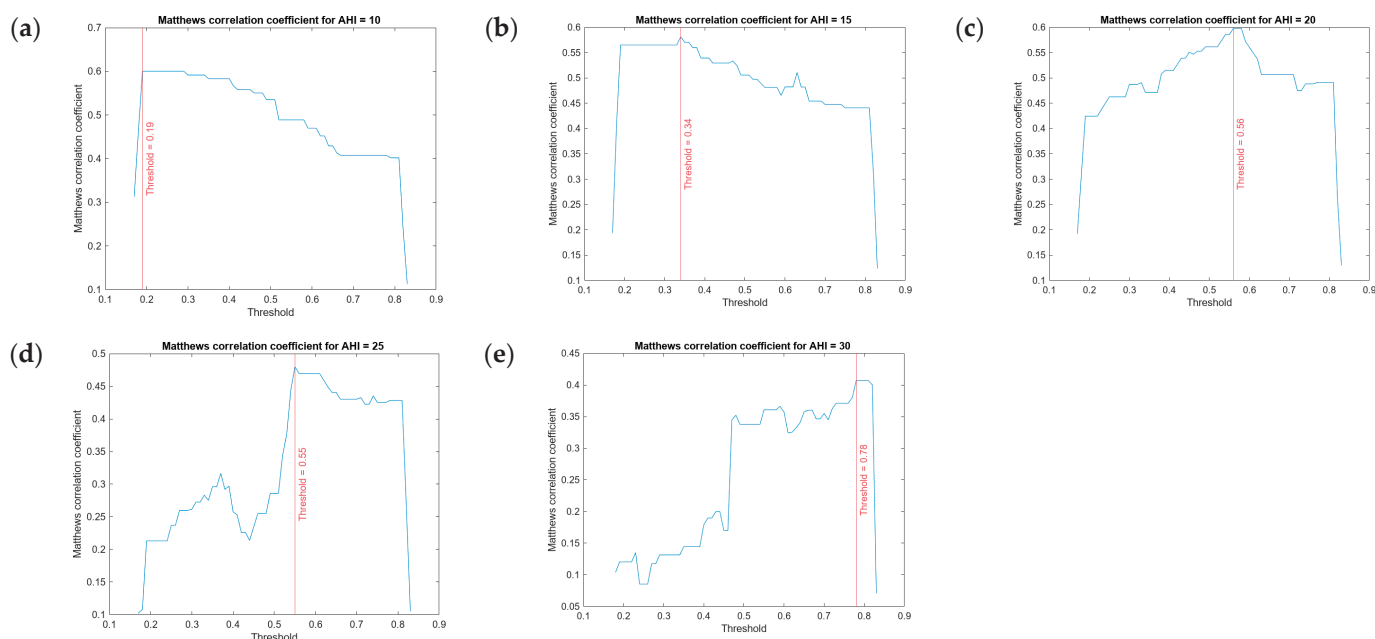


Figure 11. Determination of optimal cut points for the different inference systems: (a) AHI = 10, (b) AHI = 15, (c) AHI = 20, (d) AHI = 25, and (e) AHI = 30.

Furthermore, taking these thresholds into account, Table 9 below provides a summary of the sensitivity and specificity values obtained at each AHI level.

Table 9. Sensitivity and specificity for each AHI level at the output of the system contemplating the cut point that maximizes Mcc .

	Sensibility	Specificity
AHI 10	0.83	0.88
AHI 15	0.88	0.71
AHI 20	0.93	0.62
AHI 25	0.64	0.83
AHI 30	0.72	0.69

The above procedure makes it possible to identify an optimal point at which to transform the output values of the aggregation process into bivalued values, that is, with only two values that could be associated with suffering or not suffering from an AHI with the corresponding threshold. Since it is an iterative and non-generalized process, it should be seen as a punctual improvement of each cut-off value, trying to find that one giving the best results, analyzed as if it were a binary classification algorithm. Since it is not a generalizable procedure, it could in a way represent an additional uncertainty value since, in fact, the search is completed by results, which could lead to biased or incorrect interpretations of the cut-off. However, the sensitivity and specificity analysis, as well as the Matthews correlation coefficient itself, generally assess the accuracy of the classifier at each cut-off point, so the assumption that there is one that maximizes these values (namely the Matthews correlation coefficient) should not be seen as a fact that generates a formal explanation but as a simple measure that allows us to statistically measure the success of the classifier.

2.2.3. Generation of Alerts and Decision-Making

Finally, after all the data has been processed, the system will suggest either an ‘OSA case’ or a ‘non OSA case’ label for each AHI level. The medical team will select the AHI level that they find convenient, and then the system will generate recommendations. All this information is shown in the purple panel of Figure 8, in region 3.

3. Case Study

Once the ICDSS architecture has been introduced, this section presents a practical case that aims to demonstrate the operation of the system and highlight its potential use. This work does not intend to carry out an intensive clinical validation of the ICDSS, although an independent set of test data, not included in the dataset used for the training and construction of the model, has been reserved. It is actually a proof of concept that takes into account the results previously obtained and aims to estimate the applicability of the system.

The patient to be analyzed in the case study was extracted from this reserved test dataset.

3.1. Collection of the Patient’s Information

Table 10 shows the data of a patient suspected of suffering from OSA who was not included in the process of training and building the model and which will be analyzed in this case study.

It is important to point out that this patient underwent specific sleep tests, more specifically a cardiorespiratory polygraphy, which showed an AHI of 23.

Table 10. Data on the case study patient.

		Data	Value
Objective data		Gender	Man
		Age	69
		Weight	93 kg
		Size	179 cm
		BMI	29.03
		Neck circumference length	43 cm
		Habits	Drinking habits: daily, 30 g of alcohol
		Drug treatments	-
Subjective data	Sleep time subgroup	Hours of sleep	8 h
		Minutes until falling asleep	30 min
		Prolonged intra-sleep awakenings	No
	Unrefreshing sleep subgroup	Feeling of unrefreshing sleep	Occasionally
		Daytime tiredness	Occasionally
		Morning dullness	No
	Complicating sleep factors	Unjustified multiple awakenings	No
		Nocturia	Occasionally
		Breathless awakenings	No
	Snore subgroup	Reported apneas	No
		Snorer	Yes
		High-intensity snorer	No
	Sleep test	Snore-related awakenings	No
		AHI	23

3.2. Data Processing

Once the patient's data have been collected and entered into the application, as can be seen in Figure 12, it is proceeded to its processing by the ICDSS.

The screenshot displays the 'Apnea Diagnosis - Wang Mendel Version App' interface. It is divided into several sections:

- Initial data collection (1):** Contains 'Objective data' (General and anthropometric data, Habits, Pharmacological treatments, Illnesses) and 'Subjective data' (Sleep time - R1 a, Prolonged intra-sleep awakenings, Unrefreshing sleep - R1 b, Feeling of unrefreshing sleep, Daytime tiredness, Morning dullness, Complicating sleep factors - R2 a, Nocturia, Breathless awakenings, Reported apneas, Snore related awakenings, Snore - R2 b).
- Alert generation & Decision-making (3):** Includes a 'Generate recommendation' section with a 'Select AHI threshold' dropdown (set to 15) and a 'Calculate (4)' button. A red box indicates a 'Possible apnea case'.
- Data Processing (2):** Shows 'Statistical risks [0, 100]' with values for R10 (70), R15 (46.67), R20 (80), R25 (33.33), and R30 (43.33). It also shows 'Symbolic risk' with values for R1 a (1), R1 b (3), R2 a (3), and R2 b (4). The 'Risk aggregation [0, 1]' section shows AHI values: AHI 10 (0.8146), AHI 15 (0.4636), AHI 20 (0.8146), AHI 25 (0.5364), and AHI 30 (0.4831).
- Knowledge base for AHI = 15:** A list of 38 rules for AHI calculation based on statistical and symbolic risks.

Figure 12. Screenshot of the results obtained in the case study.

The first set of data, highlighted in blue in Table 10, is processed, as already commented, by a set of machine learning classification algorithms to determine the *Statistical Risks*. These

risks have values of 70, 46.67, 80, 33.33, and 43.44 for AHI levels of 10, 15, 20, 25, and 30, respectively, as can also be seen in Figure 12.

With regard to the second set of data, that related to the symptoms reported by the patient, highlighted in orange in Table 10, its processing is carried out by the cascade of expert systems, determining at its output the *Symbolic Risk*, which has a value of 42.86, as can be seen in Figure 12.

After that, the risk pairs (*Statistical Risk*, *Symbolic Risk*) are determined for the different AHI levels, as summarized in Table 11.

Table 11. Risk pairs for the different AHI levels.

	Statistical Risk [0, 100]	Symbolic Risk [0, 100]
AHI = 10	70	42.86
AHI = 15	46.67	42.86
AHI = 20	80	42.86
AHI = 25	33.33	42.86
AHI = 30	43.44	42.86

Risk Aggregation

These risk pairs are aggregated at each AHI level by means of a series of inference systems using the knowledge bases determined employing the automatic rule generation approach proposed by Wang and Mendel. The set of rules generated can be seen, for AHI = 15, in Figure 12 in its most basic expression of fuzzy “IF ... THEN ... ” rules. The expert medical team can review these rules for validation. This could be used to consolidate a permanent knowledge base that learns progressively with each iteration of the system.

After the inference process, the final aggregated risk value for each AHI level, the *Apnea Risk*, is obtained, as shown in Figure 12. In summary, for AHI = 10, an *Apnea Risk* value of 0.81 is obtained; for AHI = 15, an *Apnea Risk* value of 0.46 is obtained; for AHI = 20, an *Apnea Risk* value of 0.81 is obtained; for AHI = 25, an *Apnea Risk* value of 0.54 is obtained; for AHI = 30, an *Apnea Risk* value of 0.48 is obtained.

3.3. Generation of Alerts and Decision-Making

Once the *Apnea Risk* value has been determined for each AHI level, it proceeds to the last stage, in which the generation of alerts and recommendations is addressed.

In this sense, a series of colored lights are generated to facilitate the assessment of the severity of the condition, as shown in Figure 12. For this, each *Apnea Risk* value is interpreted by applying a threshold level (see Figure 11 for more information) for each AHI level. A red indicator means that the AHI level is equal to or greater than the level associated with the indicator light itself. Otherwise, the indicator will be green-colored.

In this case, it is observed that the red indicators light up for AHI levels of 10, 15, and 20, from which it is concluded that the patient could present an AHI value within the interval [20, 25). This assessment is compatible with the results that the patient obtained after carrying out the test, in which they presented a value of 23.

On the other hand, and in relation to the generation of recommendations, the medical team sets a threshold of AHI = 15. Taking this threshold into account, the system indicates that this is a potential OSA case and that sleep studies should be performed to confirm the diagnosis and, if necessary, treat the patient.

4. Discussion

The intelligent decision support system proposed in this work is undoubtedly useful and relevant in clinical practice, as already discussed in the authors’ two previous papers [1,2]. However, in this case, it is particularly important to highlight the particularities and additional benefits that the system offers over its predecessors, especially with regard to its architecture.

As a general constitutive aspect, it is worth analyzing the combination of proposals made in the second stage of the system. Both the definition of AHI thresholds and the use of two risk metrics, each associated with a different inference process, aim to extend the usefulness of the classifier. They give it, on the one hand, a larger precision in its prediction and, on the other hand, a greater representativeness of the information in the search for a solution hypothesis that generalizes such a prediction. The division into thresholds, analyzed exclusively in statistical learning, understands the model used, in this case bagged trees, as a multi-class classifier, which makes it possible to find relationships between the data that are more in line with a real prediction. Cross-validation aims to reduce the overfitting that occurs in problems of excessive complexity, an issue that is also addressed by the aggregation. The generation of these *Statistical Risks* is completed by obtaining a *Symbolic Risk* value, general and unique per patient, independent of thresholds, derived from a fuzzy inference process through a cascade of different expert systems. The initial problem of determining the risk of classifying a patient at a given threshold of the AHI can be generalized by the *Statistical* and *Symbolic Risk* pairs obtained for each of these thresholds, the statistical one being variable while the symbolic one is constant. The generalization capabilities of the intelligent system are thus limited to the adaptation of these risks according to the associated labels in the training set. If this fitting, which is easily assimilated to an aggregation, were carried out in a trivial way, with classical aggregation operators or even statistical learning models, it would not be possible to capture its significance in terms of its influence on the final diagnosis, beyond its quantitative interpretation. Therefore, in order to take advantage of the combination of the best proposals in this work, a particular aggregation approach is proposed.

Of particular relevance is therefore the third stage of the intelligent system, where, after the calculation of the risk pair (*Statistical Risk*, *Symbolic Risk*), a final value is determined for each AHI level, the *Apnea Risk*. In this sense, most of the existing approaches in the current literature usually opt for the use of analytical expressions [3,4,10,12] (aggregation operators [36–39], utility functions [2,10], curves [3,4,12], etc.), which implies knowing a priori a model that allows risks to be joined, thus introducing vagueness into the process. In general, it is possible to improve this analytical expression through stochastic optimization processes, that is, by iteratively making changes to the expression that improve the results obtained or even by proposing heuristic models adapted to the problem, which can even be solved by optimization. However, both analytical expressions and optimizable models allow aggregations to be made on the assumption that the elements to be aggregated are data, that is, point estimates of the variables of a problem. They are statistical aggregations that, as in the simplest case of a weighted sum, combine data of the same type without any further interpretation of the logic of the aggregation itself. While this is not a major concern for most problems, it could become so in medical applications, such as clinical decision support systems. In order to understand this, and as already discussed in the introduction, one must turn to the essential distinction between data and knowledge. While knowledge is a complex expression of the relationships between the qualitative and quantitative nature of the variables of a problem, data is just the volatile and quantitative expression of a variable without any greater meaning or relevance. Thus, knowledge, expressed in terms of conditional deductive rules, is a permanent construct that grows and changes with the emergence of new knowledge. Data are point values, non-permanent and changeable according to circumstances, whose value in intelligent systems is linked to their number and the ability to find plausible hypotheses of relationships through statistical learning. For this reason, aggregation has usually been applied to data, especially when it has a quantifiable expression. However, when the data represent a variable whose meaning can be inferred, as is the case with the concepts of *Statistical Risk* and *Symbolic Risk*, aggregation of these variables is not so obvious. For example, assuming a patient has values for the *Statistical Risk* of 50 and the *Symbolic Risk* of 90, what would the intelligent system suggest by aggregating the two terms? Obviously, any analytical model would process both terms as numbers and return a final numerical value depending on factors such as

the weight or importance of the two risks, a utility function linking them to qualitative assessments, a behavioral curve, or a function that is specific to the aggregation model. In any case, the number would be a statistical result that would not evaluate aspects such as the reliability of the definition of the *Symbolic Risk*, the mechanisms of the inference process, the presence or absence of knowledge in the reasoning, or in general, a non-mathematical logical explanation of the union. This is where a possible alternative to aggregation based on multivalued logics comes into play, consisting of the use of fuzzy inference approaches. This would require a fuzzy interpretation of the numerical values of the risks, assigning them to a fuzzy set, and defining the knowledge base, that is, the set of rules that define how they should be combined. In fact, this process summarizes the essential operation of an expert system in its variants based on logical rules, probability, and fuzzy rules. Expert systems are able to aggregate symbols in an explanatory way using formal logical languages and consistent bi- and multi-valued logics. They are therefore able to deal with knowledge in its most essential and basic form, as long as that knowledge can be traced and defined, either by establishing correspondences between data and symbols, or by generating rules from the data itself. However, both ways of defining knowledge are difficult, and this in part explains why statistical learning has become so relevant. In this paper, for example, both of these approaches are considered. The correspondence between data and symbols in the knowledge base is handled directly by the expert system cascade, where a set of plausible fuzzy rules has been created with the variables, which in turn have been expressed as membership functions that perform precisely this correspondence. More details on this implementation can be found in the works of the authors [1,2]. The second way, obtaining a knowledge base directly from the data, is the most relevant in this work and, as said, is its main contribution. It addresses the data-knowledge duality that exists in the genesis of any knowledge base, making it possible to consolidate knowledge from ephemeral data and offering the improvements of aggregation based on expert systems while reducing the difficulty of creating the necessary knowledge bases.

Specifically, in this work, this form of aggregation is carried out by using the method for the automatic generation of rules proposed by Wang and Mendel [40], from which the knowledge bases used by an inference system at each AHI level to determine the *Apnea Risk* value are defined. Essentially, it is a symbolic inference process that starts from numerical data obtained from the previous inference processes, that is, the machine learning classification algorithms and the cascade of expert systems.

The novelty, therefore, does not lie in the use of the inferential capacity of expert systems applied to risk aggregation but in the automatic generation of their knowledge bases, which is a notable differentiation from previous work along the same lines. By overcoming the difficulty of generating knowledge from numerical data, the use and applicability of expert systems, in this case fuzzy ones, are significantly extended.

4.1. Advantages and Disadvantages of the Aggregation Process

Given the difficulty of finding a knowledge base grounded on the experience and knowledge of the medical team, risk aggregation usually has to be carried out using statistical methods rather than the preferable symbolic methods. However, the use of the Wang-Mendel algorithm allows the knowledge base to be generated automatically, thus reducing the initial difficulty and offering the following advantages:

- By dealing with the risk values, the algorithm acts indirectly by reducing the enormous dimensionality of the problem by reducing and synthesizing the relationships between the initial variables through their representation and influence on the *Statistical* and *Symbolic Risks*.
- There is no need to define and find explicit knowledge for risk aggregation. The numerical data representing them, the risks, as well as the labels that complete each line of data can be interpreted as a non-formal pseudo-logical statement that serves as a basis for establishing a fuzzy rule.

- It is possible to combine the automatically generated rules with others generated by experts, thus completing a knowledge base with a more general scope and wider applicability.
- It allows a circumstantial logical formalization of the knowledge derived from the direct interpretation of numerical data, which is a differential milestone. This incorporates explanatory capabilities into the aggregation model and allows its validation by logical procedures.
- It allows the diversification of the risk aggregation knowledge base in its most basic form. This means that it is possible to understand that the rules form a knowledge base that can be exported to other examples once they have been created and defined.
- It reduces uncertainty in all its variants and interpretations, not only by incorporating fuzzy systems but also by not introducing aggregation models with optional or optimizable parameters.
- From the different proofs of concept elaborated, of which the case study is an example, it is possible to observe a trend towards finding accurate classifiers for each of the AHI thresholds identified. The AUCs of the ROC curves for each of them extend their area above 0.74 (1 would be a perfect classifier), which in itself would not represent a notable difference with those obtained by other simple or combined machine and deep learning algorithms if one did not take into account a differentiating factor: explainability.
- Explainability is a fair measure of the differential value of automated fuzzy rule generation from numerical data sets. Although this model, unlike traditional knowledge acquisition models, allows the aggregation of parameters with symbolic inference without the obligation of creating logical knowledge bases, it shares with them the ability to explain—that is, to understand, comprehend, and reason—the inference process. However, in order for these new bases to have a permanent, not to say ontological, character, they must be verified by the team of expert pulmonologists.

Similarly, although its advantages are evident from the results of the conceptual tests of the developed system, there are also some disadvantages that need to be considered:

- The first origin of the rules is the set of objective and subjective, numeric and categorical data from which the initial set of objective and subjective data, represented by its risk measures and class labels, was derived. These data are, by definition, ephemeral and mutable, so the rules derived from them will inherit these characteristics. Therefore, although the result of the application of the algorithm is a set of fuzzy rules, they will necessarily need to be interpreted by the expert team in order to consolidate them into a permanent knowledge base.
- Although the application of the algorithm results in a set of fuzzy rules, it is essential that they be interpreted by a team of experts in order to consolidate them into a permanent knowledge base.
- The logical operators that relate the antecedents must be defined in advance, which implies a prior assumption that may not be true.
- The algorithm classifies and ranks the rules according to a degree value, calculated as the product of the membership values of the set of antecedents and consequents of each rule. In the case of rules having similar antecedents, those with a higher degree are given priority, and the others are eliminated. This, in fact, implies a loss of knowledge due to the operation of the algorithm itself, which will have to be taken into account in due course. In more or less well-known applications where prior knowledge exists, this is not a very significant loss. However, in medical diagnosis it may lead to non-negligible error rates, and therefore, in the future, it will be necessary to study mechanisms to reduce the loss of knowledge generated by the approach of generating deductive rules on fuzzy sets proposed by Wang and Mendel [40].

4.2. Comparison of Systems

In order to better understand the different contributions of the current proposal, in Table 12, a comparison of the current system with other similar systems is proposed. The criteria used are as follows: ease of use, related to the difficulty for a team of medical experts

to prepare the precise conditions for calculating the system's prediction; knowledge acquisition, related to the existence within the intelligent system of an explicit and automated knowledge acquisition subsystem, distinguished from data; data dependency, in that the system has a structural dependency on a statistical learning model; and combination of inference models, present when the intelligent system effectively combines, hybridized or not, different inference mechanisms of a heterogeneous nature.

Table 12. Comparison of the current system with other similar systems.

	Ease of Use (Easy/Moderate/ Complex)	Automated Knowledge Acquisition (Full/Partial)	Data Dependency (Full/Partial)	Combination of Inference Models (Yes/No)
Ismail Atacak [65]: This paper proposes a malware detection system for the Android operating system. To this end, six classification machine learning algorithms are first used, focusing on the processing of application data. Then, through a voting process, three of the algorithms are selected, whose results are interpreted and aggregated by a Mamdani-type fuzzy inference system, which makes it possible to determine the degree of malware.	Moderate difficulty of use. To use this system, it is essential to define in advance the knowledge base of the fuzzy inference system that will allow the aggregation of the results of the machine learning algorithms.	Partial knowledge acquisition. The system architecture does not have an automatic knowledge acquisition sub-system. The rules of the Mamdani inference system need to be defined manually.	Full data dependency. This is a data dependent approach as it uses supervised machine learning algorithms.	Yes, it consists of the sequential use of a block of machine learning algorithms and a fuzzy inference system, which can be understood as an ensemble, that is, where the fuzzy engine allows the predictions of the preceding algorithms to be aggregated.
	-	-	=	=
Melin et al. [66]: In this work, aimed at predicting the COVID-19 time series, it is proposed to use jointly a block of neural networks, more specifically nonlinear autoregressive neural networks and function fitting neural networks, and a fuzzy inference system focused on determining the importance of the predictions of the models, which are aggregated through a weighted summation model.	Moderate difficulty of use. To use it, it is essential to first define the knowledge base that will allow the importance of each model's predictions to be determined on the basis of its prediction error.	Partial knowledge acquisition. The system integrates a module based on the Mamdani inference system. Its rules are defined manually, as it does not include a knowledge acquisition subsystem.	Full data dependency. This is an unavoidably data-dependent approach where it is necessary to train the neural networks used.	Yes, this paper uses together both statistical inferential approaches and symbolic inferential approaches.
	-	-	=	=
Ahmed et al. [67]: This paper proposes a system for the prediction of diabetes condition. To this end, two machine learning algorithms are used that focus on the processing of initial data, whose predictions are then processed by an inference system based on fuzzy logic of the Mamdani type, which is responsible for determining the final prediction.	Moderate difficulty of use. In order to use this system, it is necessary to first define the knowledge base of the Mamdani inference system, which is essential in determining the prediction of the model.	Partial knowledge acquisition. The knowledge base is manually defined, as the system lacks a knowledge acquisition subsystem.	Full data dependency. As it integrates approaches based on the use of classification machine learning algorithms, the availability of data is essential.	Yes, this system uses both statistical and symbolic inferential approaches in a joint and sequential manner.
	-	-	=	=
Ragman et al. [68]: This paper proposes a real-time rainfall forecasting system. To achieve this, it uses four classification machine learning algorithms that focus on the processing of sensor data. The predictions of these models are combined by a Mamdani inference system, which is responsible for determining the final prediction.	Moderately difficult to use. In order to use the system, it is essential to define the knowledge base of the fuzzy inference system responsible for aggregating the predictions.	Partial knowledge acquisition. The system does not have a knowledge acquisition subsystem. The knowledge base is manually defined.	Full data dependency. This is a data dependent approach as it incorporates supervised learning approaches into its architecture.	Yes. It uses statistical and symbolic inference approaches.
	-	-	=	=

Table 12. Cont.

	Ease of Use (Easy/Moderate/ Complex)	Automated Knowledge Acquisition (Full/Partial)	Data Dependency (Full/Partial)	Combination of Inference Models (Yes/No)
Casal-Guisande et al. [2]: This is an intelligent decision support system applied to the diagnosis of obstructive sleep apnea. Its architecture uses a set of classification machine learning algorithms and a cascade of expert systems, each of which outputs a risk indicator. These indicators are then combined through a utility function that determines a metric associated with suffering from the pathology.	Moderate difficulty of use. In order to use the intelligent system, it is necessary to have defined the knowledge bases of the different expert systems. In this sense, after this milestone, the system could be used without major difficulties other than those related to the revision and improvement of the rules for its use.	Partial knowledge acquisition. The architecture of the intelligent system, more specifically the cascade of expert systems, does not include a specific knowledge acquisition subsystem. This must be manually defined by the expert team.	Full data dependency. The data is required to train the machine learning classification algorithms.	Yes. It uses symbolic and statistical inference approaches in a concurrent mode. Their results are then combined using a specific utility function.
	-	-	=	=
Casal-Guisande et al. [5]: This is an intelligent decision support system applied to the diagnosis of breast cancer, focusing on the interpretation of information obtained from mammograms. Its sequential architecture uses a cascade of expert systems, whose output is a set of risk indicators. A set of underlying factors that summarise and represent the risks is then obtained by applying factor analysis approaches. These are processed by a classification machine learning algorithm, which makes it possible to determine a risk metric associated with suffering from the pathology.	Moderate difficulty of use. The use of the intelligent system requires the definition of the knowledge bases of the different expert systems. In this sense, after this milestone, the system could be used without major difficulties other than those related to the revision and improvement of the rules.	Partial knowledge acquisition. The architecture of the intelligent system, and more specifically the cascade of expert systems, does not contemplate a specific subsystem for acquiring knowledge. The knowledge bases must be defined manually by the expert team.	Full data dependency. Due to its sequential nature and the use of a machine learning classification algorithm, this approach requires a data set from the beginning.	Yes. It uses symbolic and statistical inferential approaches.
	-	-	=	=
Casal-Guisande et al. [10]: This paper presents an intelligent system applied to the diagnosis of breast cancer, focusing on the interpretation of the information obtained after performing mammograms. Its architecture uses a set of expert systems and a classification machine learning algorithm working concurrently. The output is a set of risk indicators that are combined using a specific analytical function to obtain a risk indicator. In addition, the system incorporates a corrective approach that allows the weighting of the risk obtained through the opinions of the experts, which is reflected in the BI-RADS indicator.	Moderate difficulty of use. For the use of the intelligent system, it is necessary to define the knowledge bases of the different expert systems. In this sense, after this milestone, the system could be used without major difficulties other than those associated with the revision and improvement of the rules.	Partial knowledge acquisition. There is no subsystem for knowledge acquisition. The knowledge bases of expert systems must be defined manually.	Full data dependency. The system is data dependent due to the use of a classification machine learning algorithm.	Yes, the system integrates various inferential approaches, both symbolic and statistic.
	-	-	=	=

Table 12. Cont.

	Ease of Use (Easy/Moderate/ Complex)	Automated Knowledge Acquisition (Full/Partial)	Data Dependency (Full/Partial)	Combination of Inference Models (Yes/No)
Our proposal	Very easy to use. Beyond the definition of the knowledge bases of the cascade of expert systems, it is not necessary to define the rules of the Mamdani inference system responsible for risk aggregation. The system can be used from the beginning without major difficulties.	Full and automatic knowledge acquisition. The system integrates an automatic fuzzy rule generation mechanism, the algorithm proposed by Wang and Mendel.	Full data dependency. In addition to defining the classification machine learning algorithms, data is also needed to create the corpus of rules for the aggregation inference system.	Yes, the system integrates a variety of inference approaches, both symbolic and statistical.

5. Conclusions

In healthcare environments, the use of intelligent systems to support medical teams in diagnostic processes is becoming increasingly common. In this sense, this work addresses the improvement and evolution of an intelligent decision support system for the diagnosis of OSA cases. For this purpose, combining proposals previously published by the authors and starting from patient information, a series of machine learning classification algorithms are used, as well as a series of expert systems arranged in cascade, with the aim of obtaining a series of risk pairs (*Statistical Risk*, *Symbolic Risk*), each focused on an AHI level. Each risk pair is then processed by a subsequent inference system whose knowledge base is automatically generated. For this task, an automatic rule generation approach is used, specifically the one proposed by Wang and Mendel [40], which makes it possible to determine the *Apnea Risk* value for each AHI level.

The intelligent system has been implemented as a software artifact, and its operation has been demonstrated by means of a case study, which has made it possible to highlight the usefulness of the system as a tool to support the diagnostic process. It should also be noted that the tests carried out on a test dataset with 183 patients, independent of those used to construct the model, showed AUC values between 0.74 and 0.88 and Matthews correlation coefficient values between 0.41 and 0.6 for the different AHI levels.

The inclusion of automatic rule generation approaches in the architecture of intelligent systems opens up several promising lines for future research and development. Undoubtedly, the main one lies in the effective integration of this type of rule generator into the architecture of expert systems, in line with the work already pointed out in the genesis of second-generation expert systems. Similarly, their use in the new generation of hybrid intelligent systems needs to be explored and studied in detail. Likewise, the validation of these knowledge bases, with a view to their transformation into permanent ontological bases, is also an unresolved issue, especially from a point of view that is not fully assisted by the human expert. Furthermore, the Wang-Mendel algorithm itself can be redefined to avoid the inherent loss of information associated with eliminating low-degree rules by proposing a way to integrate this information into new high-degree rules. It can even be revised and improved with different logical operators that relate antecedents. There are undoubtedly many ways to improve this promising approach to knowledge acquisition.

Author Contributions: Conceptualization, M.C.-G. and A.C.-C.; methodology, M.C.-G. and A.C.-C.; software, M.C.-G.; validation, M.C.-G., A.C.-C., J.-B.B.-R. and J.C.-P.; investigation, M.C.-G., A.C.-C., J.-B.B.-R. and J.C.-P.; data curation, M.C.-G.; writing—original draft preparation, M.C.-G.; writing—review and editing, A.C.-C. and J.C.-P.; supervision, J.-B.B.-R. and J.C.-P. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Institutional Review Board Statement: The study was conducted in accordance with the Declaration of Helsinki, and approved by the Ethics Committee of Galicia (protocol code 2022/256, 2 July 2022).

Data Availability Statement: Not applicable.

Acknowledgments: M.C.-G. is grateful to *Consellería de Educación, Universidade e Formación Profesional e Consellería de Economía, Emprego e Industria da Xunta de Galicia* (ED481A-2020/038) for his pre-doctoral fellowship.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Casal-Guisande, M.; Torres-Durán, M.; Mosteiro-Añón, M.; Cerqueiro-Pequeno, J.; Bouza-Rodríguez, J.-B.; Fernández-Villar, A.; Comesaña-Campos, A. Design and Conceptual Proposal of an Intelligent Clinical Decision Support System for the Diagnosis of Suspicious Obstructive Sleep Apnea Patients from Health Profile. *Int. J. Environ. Res. Public Health* **2023**, *20*, 3627. [CrossRef] [PubMed]
2. Casal-Guisande, M.; Ceide-Sandoval, L.; Mosteiro-Añón, M.; Torres-Durán, M.; Cerqueiro-Pequeno, J.; Bouza-Rodríguez, J.-B.; Fernández-Villar, A.; Comesaña-Campos, A. Design of an Intelligent Decision Support System Applied to the Diagnosis of Obstructive Sleep Apnea. *Diagnostics* **2023**, *13*, 1854. [CrossRef]
3. Casal-Guisande, M.; Bouza-Rodríguez, J.B.; Comesaña-Campos, A.; Cerqueiro-Pequeno, J. Proposal and Definition of a Methodology for Remote Detection and Prevention of Hypoxemic Clinical Cases in Patients Susceptible to Respiratory Diseases. In Proceedings of the Seventh International Conference on Technological Ecosystems for Enhancing Multiculturality, León, Spain, 16–18 October 2019; pp. 331–338. [CrossRef]
4. Casal-Guisande, M.; Comesaña-Campos, A.; Cerqueiro-Pequeno, J.; Bouza-Rodríguez, J.-B. Design and Development of a Methodology Based on Expert Systems, Applied to the Treatment of Pressure Ulcers. *Diagnostics* **2020**, *10*, 614. [CrossRef] [PubMed]
5. Casal-Guisande, M.; Comesaña-Campos, A.; Dutra, I.; Cerqueiro-Pequeno, J.; Bouza-Rodríguez, J.-B. Design and Development of an Intelligent Clinical Decision Support System Applied to the Evaluation of Breast Cancer Risk. *J. Pers. Med.* **2022**, *12*, 169. [CrossRef] [PubMed]
6. Casal-Guisande, M.; Cerqueiro-Pequeno, J.; Comesaña-Campos, A.; Bouza-Rodríguez, J.B. Proposal of a Methodology Based on Expert Systems for the Treatment of Diabetic Foot Condition. In Proceedings of the Eighth International Conference on Technological Ecosystems for Enhancing Multiculturality, Salamanca, Spain, 21–23 October 2020; ACM: New York, NY, USA; pp. 491–495.
7. Casal-Guisande, M.; Comesaña-Campos, A.; Pereira, A.; Bouza-Rodríguez, J.-B.; Cerqueiro-Pequeno, J. A Decision-Making Methodology Based on Expert Systems Applied to Machining Tools Condition Monitoring. *Mathematics* **2022**, *10*, 520. [CrossRef]
8. Cerqueiro-Pequeno, J.; Comesaña-Campos, A.; Casal-Guisande, M.; Bouza-Rodríguez, J.-B. Design and Development of a New Methodology Based on Expert Systems Applied to the Prevention of Indoor Radon Gas Exposition Risks. *Int. J. Environ. Res. Public Health* **2020**, *18*, 269. [CrossRef]
9. Cerqueiro-Pequeno, J.; Casal-Guisande, M.; Comesaña-Campos, A.; Bouza-Rodríguez, J.-B. Conceptual Design of a New Methodology Based on Intelligent Systems Applied to the Determination of the User Experience in Ambulances. In Proceedings of the Ninth International Conference on Technological Ecosystems for Enhancing Multiculturality (TEEM'21), Barcelona, Spain, 26–29 October 2021; pp. 290–296. [CrossRef]
10. Casal-Guisande, M.; Álvarez-Pazó, A.; Cerqueiro-Pequeno, J.; Bouza-Rodríguez, J.-B.; Peláez-Lourido, G.; Comesaña-Campos, A. Proposal and Definition of an Intelligent Clinical Decision Support System Applied to the Screening and Early Diagnosis of Breast Cancer. *Cancers* **2023**, *15*, 1711. [CrossRef]
11. Casal-Guisande, M.; Bouza-Rodríguez, J.-B.; Cerqueiro-Pequeno, J.; Comesaña-Campos, A. Design and Conceptual Development of a Novel Hybrid Intelligent Decision Support System Applied towards the Prevention and Early Detection of Forest Fires. *Forests* **2023**, *14*, 172. [CrossRef]
12. Comesaña-Campos, A.; Casal-Guisande, M.; Cerqueiro-Pequeno, J.; Bouza-Rodríguez, J.B. A Methodology Based on Expert Systems for the Early Detection and Prevention of Hypoxemic Clinical Cases. *Int. J. Environ. Res. Public Health* **2020**, *17*, 8644. [CrossRef]
13. Suthaharan, S. *Machine Learning Models and Algorithms for Big Data Classification*; Springer: Boston, MA, USA, 2016; Volume 36, ISBN 978-1-4899-7640-6.
14. Zhou, Z.-H. *Machine Learning*; Springer: Singapore, 2021; ISBN 978-981-15-1966-6.
15. Deo, R.C. Machine Learning in Medicine. *Circulation* **2015**, *132*, 1920–1930. [CrossRef]
16. Tang, Y.; Huang, J.; Pedrycz, W.; Li, B.; Ren, F. A Fuzzy Clustering Validity Index Induced by Triple Center Relation. *IEEE Trans. Cybern.* **2023**, 1–13. [CrossRef] [PubMed]
17. Celebi, M.E.; Aydin, K. *Unsupervised Learning Algorithms*; Springer International Publishing: Cham, Switzerland, 2016; ISBN 978-3-319-24209-5.

18. Raza, K.; Singh, N.K. A Tour of Unsupervised Deep Learning for Medical Image Analysis. *Curr. Med. Imaging* **2021**, *17*, 1059–1077. [CrossRef]
19. Jabbar, H.K.; Zaman Khan, R. Survey on Development of Expert System from 2010 to 2015. In Proceedings of the Second International Conference on Information and Communication Technology for Competitive Strategies, Udaipur, India, 4–5 March 2016. [CrossRef]
20. Sheikhtaheri, A.; Sadoughi, F.; Hashemi Dehaghi, Z. Developing and Using Expert Systems and Neural Networks in Medicine: A Review on Benefits and Challenges. *J. Med. Syst.* **2014**, *38*, 1–6. [CrossRef] [PubMed]
21. Arsene, O.; Dumitrache, I.; Mihiu, I. Expert System for Medicine Diagnosis Using Software Agents. *Expert Syst. Appl.* **2015**, *42*, 1825–1834. [CrossRef]
22. Masri, R.Y.; Mat Jani, H. Employing Artificial Intelligence Techniques in Mental Health Diagnostic Expert System. In Proceedings of the 2012 International Conference on Computer & Information Science (ICCIS), Kuala Lumpur, Malaysia, 12–14 June 2012; Volume 1, pp. 495–499. [CrossRef]
23. Djam, X.Y.; Wajiga, G.M.; Kimbi, Y.H.; Blamah, N.V. A Fuzzy Expert System for the Management of Malaria. *Int. J. Pure Appl. Sci. Technol.* **2011**, *5*, 84–108.
24. Yasnoff, W.A.; Miller, P.L. *Decision Support and Expert Systems in Public Health*; Springer: London, UK, 2014; pp. 449–467.
25. Torshizi, A.D.; Zarandi, M.H.F.; Torshizi, G.D.; Eghbali, K. A Hybrid Fuzzy-Ontology Based Intelligent System to Determine Level of Severity and Treatment Recommendation for Benign Prostatic Hyperplasia. *Comput. Methods Programs Biomed.* **2014**, *113*, 301–313. [CrossRef]
26. Benjafield, A.V.; Ayas, N.T.; Eastwood, P.R.; Heinzer, R.; Ip, M.S.M.; Morrell, M.J.; Nunez, C.M.; Patel, S.R.; Penzel, T.; Pépin, J.L.D.; et al. Estimation of the Global Prevalence and Burden of Obstructive Sleep Apnoea: A Literature-Based Analysis. *Lancet Respir. Med.* **2019**, *7*, 687–698. [CrossRef]
27. Thurnheer, R.; Bloch, K.E.; Laube, I.; Gugger, M.; Heitz, M. Respiratory Polygraphy in Sleep Apnoea Diagnosis. *Swiss Med. Wkly.* **2007**, *137*, 97–102.
28. Alonso Álvarez, M.d.I.L.; Terán Santos, J.; Cordero Guevara, J.; Martínez, M.G.; Rodríguez Pascual, L.; Viejo Bañuelos, J.L.; Marañón Cabello, A. Fiabilidad de La Poligrafía Respiratoria Domiciliaria Para El Diagnóstico Del Síndrome de Apneas-Hipopneas Durante El Sueño. Análisis de Costes. *Arch. Bronconeumol.* **2008**, *44*, 22–28. [CrossRef] [PubMed]
29. Calleja, J.M.; Esnaola, S.; Rubio, R.; Durán, J. Comparison of a Cardiorespiratory Device versus Polysomnography for Diagnosis of Sleep Apnoea. *Eur. Respir. J.* **2002**, *20*, 1505–1510. [CrossRef]
30. Douglas, N.J.; Thomas, S.; Jan, M.A. Clinical Value of Polysomnography. *Lancet* **1992**, *339*, 347–350. [CrossRef] [PubMed]
31. Rundo, J.V.; Downey, R. Polysomnography. In *Handbook of Clinical Neurology*; Elsevier: Amsterdam, The Netherlands, 2019; Volume 160, pp. 381–392.
32. Kapur, V.K.; Auckley, D.H.; Chowdhuri, S.; Kuhlmann, D.C.; Mehra, R.; Ramar, K.; Harrod, C.G. Clinical Practice Guideline for Diagnostic Testing for Adult Obstructive Sleep Apnea: An American Academy of Sleep Medicine Clinical Practice Guideline. *J. Clin. Sleep Med.* **2017**, *13*, 479–504. [CrossRef] [PubMed]
33. Punjabi, N.M. The Epidemiology of Adult Obstructive Sleep Apnea. *Proc. Am. Thorac. Soc.* **2008**, *5*, 136–143. [CrossRef] [PubMed]
34. Ramachandran, A.; Karuppiyah, A. A Survey on Recent Advances in Machine Learning Based Sleep Apnea Detection Systems. *Healthcare* **2021**, *9*, 914. [CrossRef] [PubMed]
35. Mostafa, S.S.; Mendonça, F.; Ravelo-García, A.G.; Morgado-Dias, F. A Systematic Review of Detecting Sleep Apnea Using Deep Learning. *Sensors* **2019**, *19*, 4934. [CrossRef]
36. Denis Bouyssou, T.; Marchant, M.P.; Tsoukiàs, P.V. *Evaluation and Decision Models with Multiple Criteria: Stepping Stones for the Analyst*; Springer: New York, NY, USA, 2006; ISBN 978-1-4419-4053-7.
37. Theil, H. Alternative Approaches to the Aggregation Problem. *Stud. Log. Found. Math.* **1966**, *44*, 507–527. [CrossRef]
38. Pomerol, J.-C.; Barba-Romero, S. *Multicriterion Decision in Management. Principles and Practice*; Springer Science & Business: New York, NY, USA, 2000.
39. Allen, B. On the Aggregation of Preferences in Engineering Design. In Proceedings of the International Design Engineering Technical Conferences and Computers and Information in Engineering Conference, Pittsburgh, PA, USA, 9–12 September 2001; Volume 2, pp. 57–67. [CrossRef]
40. Wang, L.X.; Mendel, J.M. Generating Fuzzy Rules by Learning from Examples. *IEEE Trans. Syst. Man Cybern.* **1992**, *22*, 1414–1427. [CrossRef]
41. Mitchell, T.M. *Machine Learning*; McGraw-Hill Education: Maidenhead, UK, 1997.
42. Neale, I.M. First Generation Expert Systems: A Review of Knowledge Acquisition Methodologies. *Knowl. Eng. Rev.* **1988**, *3*, 105–145. [CrossRef]
43. Steels, L. Second Generation Expert Systems. *Future Gener. Comput. Syst.* **1985**, *1*, 213–221. [CrossRef]
44. Rubin, S.H.; Murthy, J.; Ceruti, M.G.; Milanova, M.; Ziegler, Z.C.; Rush, R.J., Jr. Third-Generation Expert Systems. In Proceedings of the 2nd International ISCA Conference on Information Reuse and Integration (IRI-2000), Honolulu, HI, USA, 1–3 November 2000; pp. 27–34.
45. Mamdani, E.H. Advances in the Linguistic Synthesis of Fuzzy Controllers. *Int. J. Man-Mach. Stud.* **1976**, *8*, 669–678. [CrossRef]
46. Mamdani, E.H. Application of Fuzzy Logic to Approximate Reasoning Using Linguistic Synthesis. *IEEE Trans. Comput.* **1977**, *C-26*, 1182–1191. [CrossRef]

47. Mamdani, E.H.; Assilian, S. An Experiment in Linguistic Synthesis with a Fuzzy Logic Controller. *Int. J. Man-Mach. Stud.* **1975**, *7*, 1–13. [CrossRef]
48. Ross, T.J. *Fuzzy Logic with Engineering Applications*, 3rd ed.; John Wiley & Sons, Ltd.: Chichester, UK, 2010; ISBN 9781119994374.
49. Boughorbel, S.; Jarray, F.; El-Anbari, M. Optimal Classifier for Imbalanced Data Using Matthews Correlation Coefficient Metric. *PLoS ONE* **2017**, *12*, e0177678. [CrossRef] [PubMed]
50. Chicco, D.; Jurman, G. The Advantages of the Matthews Correlation Coefficient (MCC) over F1 Score and Accuracy in Binary Classification Evaluation. *BMC Genom.* **2020**, *21*, 1–13. [CrossRef] [PubMed]
51. Guilford, J.P. *Psychometric Methods*; McGraw-Hill: New York, NY, USA, 1954.
52. Hevner, A.R.; Chatterjee, S. *Design Research in Information Systems: Theory and Practice*; Springer: New York, NY, USA, 2010; ISBN 978-1-4419-6107-5.
53. Hevner, A.R.; March, S.T.; Park, J.; Ram, S. Design Science in Information Systems Research. *MISQ* **2004**, *28*, 75–105. [CrossRef]
54. MATLAB App Designer—MATLAB & Simulink. Available online: <https://es.mathworks.com/products/matlab/app-designer.html> (accessed on 10 August 2022).
55. Classification Learner. Available online: <https://www.mathworks.com/help/stats/classificationlearner-app.html> (accessed on 18 October 2022).
56. Fuzzy Logic Toolbox—MATLAB. Available online: <https://www.mathworks.com/products/fuzzy-logic.html> (accessed on 1 November 2022).
57. Imbalanced-Learn. Available online: <https://imbalanced-learn.org/dev/index.html> (accessed on 18 October 2022).
58. Agresti, A. *Categorical Data Analysis*; John Wiley & Sons, Inc.: Hoboken, NJ, USA, 2002; ISBN 0471360937.
59. Powers, D.; Xie, Y. *Statistical Methods for Categorical Data Analysis*; Emerald Group Publishing: Bingley, UK, 2008.
60. Mohammed, A.J.; Hassan, M.M.; Kadir, D.H. Improving Classification Performance for a Novel Imbalanced Medical Dataset Using Smote Method. *Int. J. Adv. Trends Comput. Sci. Eng.* **2020**, *9*, 3161–3172. [CrossRef]
61. Chawla, N.V.; Bowyer, K.W.; Hall, L.O.; Kegelmeyer, W.P. SMOTE: Synthetic Minority Over-Sampling Technique. *J. Artif. Intell. Res.* **2002**, *16*, 321–357. [CrossRef]
62. Miller, G.A. The Magical Number Seven, plus or Minus Two: Some Limits on Our Capacity for Processing Information. *Psychol. Rev.* **1956**, *63*, 81–97. [CrossRef]
63. Kumari, R.; Srivastava, S. Machine Learning: A Review on Binary Classification. *Int. J. Comput. Appl.* **2017**, *160*, 11–15. [CrossRef]
64. Duda, R.O.; Hart, P.E.; Stork, D.G. *Pattern Classification*; Wiley: Hoboken, NJ, USA, 2000.
65. Atacak, I. An Ensemble Approach Based on Fuzzy Logic Using Machine Learning Classifiers for Android Malware Detection. *Appl. Sci.* **2023**, *13*, 1484. [CrossRef]
66. Melin, P.; Monica, J.C.; Sanchez, D.; Castillo, O. Multiple Ensemble Neural Network Models with Fuzzy Response Aggregation for Predicting COVID-19 Time Series: The Case of Mexico. *Healthcare* **2020**, *8*, 181. [CrossRef] [PubMed]
67. Ahmed, U.; Issa, G.F.; Khan, M.A.; Aftab, S.; Khan, M.F.; Said, R.A.T.; Ghazal, T.M.; Ahmad, M. Prediction of Diabetes Empowered With Fused Machine Learning. *IEEE Access* **2022**, *10*, 8529–8538. [CrossRef]
68. Rahman, A.U.; Abbas, S.; Gollapalli, M.; Ahmed, R.; Aftab, S.; Ahmad, M.; Khan, M.A.; Mosavi, A. Rainfall Prediction System Using Machine Learning Fusion for Smart Cities. *Sensors* **2022**, *22*, 3504. [CrossRef] [PubMed]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

Article

Comparative Study of Type-1 and Interval Type-2 Fuzzy Logic Systems in Parameter Adaptation for the Fuzzy Discrete Mycorrhiza Optimization Algorithm

Hector Carreon-Ortiz, Fevrier Valdez and Oscar Castillo *

Tijuana Institute of Technology, TecNM, Tijuana 22379, Mexico; hector.co@tectijuana.edu.mx (H.C.-O.); fevrier@tectijuana.mx (F.V.)

* Correspondence: ocastillo@tectijuana.mx

Abstract: The Fuzzy Discrete Mycorrhiza Optimization (FDMOA) Algorithm is a new hybrid optimization method using the Discrete Mycorrhiza Optimization Algorithm (DMOA) in combination with type-1 or interval type-2 fuzzy logic system. In this new research, when using T1FLS, membership functions are defined by type-1 fuzzy sets, which allows for a more flexible and natural representation of uncertain and imprecise data. This approach has been successfully applied to several optimization problems, such as in feature selection, image segmentation, and data clustering. On the other hand, when DMOA is using IT2FLS, membership functions are represented by interval type-2 fuzzy sets, which allows for a more robust and accurate representation of uncertainty. This approach has been shown to handle higher levels of uncertainty and noise in the input data and has been successfully applied to various optimization problems, including control systems, pattern recognition, and decision-making. Both DMOA using T1FLS and DMOA using IT2FLS have shown better performance than the original DMOA algorithm in many applications. The combination of DMOA with fuzzy logic systems provides a powerful and flexible optimization framework that can be adapted to various problem domains. In addition, these techniques have the potential to more efficiently and effectively solve real-world problems.

Keywords: discrete; optimization; type-2 fuzzy logic system; metaheuristic

MSC: 03B52; 03E72; 62P30

1. Introduction

Heuristic methods involve searching for solutions through trial and error, while metaheuristics are considered more advanced because they use information and solution selection to guide the search process [1]. Most metaheuristics imitate nature, specifically biological systems that have evolved over time due to natural selection [2]. For instance, ticks rely on temperature and body odor as important indicators, while bats use air compression waves to echo in caves. These unique traits have inspired algorithms that imitate nature and have gained popularity in various fields, such as fuzzy systems, neural networks, machine learning, artificial intelligence, computational intelligence, and engineering. These applications often require sophisticated optimization algorithms due to their involvement in nonlinear optimization [1,3,4].

Algorithms are used to solve optimization problems, but uncertainty in the real world can make this search more complicated. To address this, an optimal and robust design is aimed to find the best possible solutions. Solutions that are optimal but not robust are not practical in the real world [1,5–7].

This study proposes a new optimization algorithm called the Mycorrhiza Optimization Algorithm (MOA), which is inspired by the symbiotic relationship between plant roots and fungi shown in Figure 1. The algorithm aims to optimize resource allocation

and communication through biochemical signals that alert the organisms to predators or other dangers.

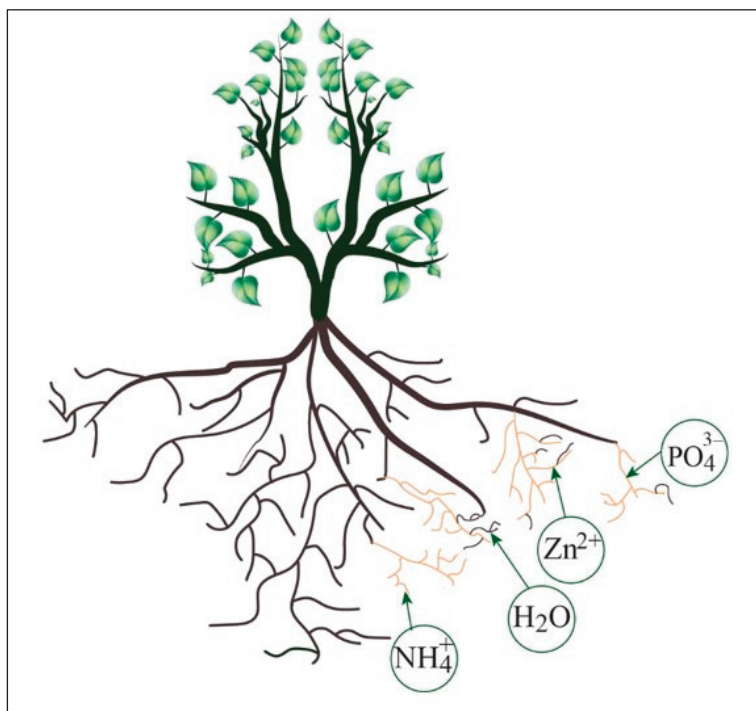


Figure 1. Symbiosis between plant and MN.

Maximizing or minimizing is a technique that has been used for a long time to achieve optimized results. It is a common practice in everyday life and can be applied to various areas such as time, money, and resources. As a result, optimization is becoming increasingly important [8].

Symbiotic fungi have been thought to help plants migrate to land by improving the search for mineral nutrients and exchanging them for photosynthetic organic carbon. Nowadays, plant–fungal symbioses are common and varied. Recent findings suggest that early terrestrial plants had access to a range of possible fungal associations and that these associations were influenced by changes in atmospheric CO_2 concentrations.

The relationship between soil-dwelling filamentous fungi and plants is one of the most significant examples of natural symbiosis, following mitochondria and plastids. Most land-based plants, including many agricultural crops, form close symbiotic relationships with fungi. These mutually beneficial partnerships are believed to have played a vital role in the evolution and diversity of land plants. In plants with roots, this relationship is known as “mycorrhiza,” while in plants without roots that have intracellular fungal structures resembling those in rooted plants, such as coils and arbuscules, it is called “mycorrhiza-like.”

Fungal mycelium can spread extensively through soil, and the relationships between fungi and host plants are often widespread, leading to the development of Mycorrhiza Networks (MN). These networks consist of uninterrupted fungal mycelia that connect two or more plants, regardless of whether they are of the same or different species. This allows MN to integrate numerous plant and fungal species, which interact, give feedback, and adapt, creating a complex adaptive social network. It is becoming increasingly clear that MN has an impact on member plants and fungi and involves communication between plants and fungi through biochemical signaling, resource transfers, or the influence of electrical signals.

Plants and fungi respond quickly to communication through MN, which can be described as behavioral responses. This approach allows us to view MN in terms of plant behavior.

A design approach for the algorithm involves five phases that incorporate a simulation of the ecosystem that exists between plants and the mycorrhizal fungal network. The five phases are colonization, mutualism, pollination, resource exchange, and plant defense. The Lotka–Volterra system of equations is utilized in three of these phases, which involve discrete equations. In the pollination phase, Levy’s flight equation can be used, but other equations can also be used.

The main contribution of this research is enhancing the ability of the Discrete Mycorrhizal Optimization Algorithm (DMOA) in conjunction with the type-1 fuzzy logic system (T1FLS) and interval type-2 fuzzy logic system (IT2FLS) to obtain better results than those found by only using the original algorithm and how finally the results of the DMOA together with the IT2FLS were the best of all, which is what we expected, due to the ability of IT2FLSs of handling higher degrees of uncertainty

The article is organized as follows, in Section 1 we make an introduction to the Discrete Mycorrhizal Optimization Algorithm, in Section 2 we make a description of the concept of optimization and what it actually represents for an algorithm, Section 3 shows the basics of the T1FLS and IT2FLS, Section 4 shows a brief history of the developed algorithms in relation to the plants, in Section 5 we present the method with which we developed this investigation, in Section 6 we present the results by means of tables and graphs, and in Section 7 we present the conclusions of this investigation.

2. Optimization

The term “optimization” refers to the process of selecting the best possible option among several viable options while keeping a set of constraints in mind. Optimization theory is an example of how humans seek perfection by teaching how to define and achieve an optimal outcome. Through optimization, the aim is to enhance system performance towards the optimal point(s) [9]. Depending on the theoretical or practical aspect, optimization can be considered a component of applied or numerical mathematics or a computer-based system design method [10]. The analytical solution of an optimization problem relies on the shape of the criterion and constraint functions [11]. The simplest form of the general optimization problem is the unconstrained optimization issue, where decision variables have no restrictions, and calculus can be applied for analysis. Another relatively simple form is when all the constraints can be expressed as equality relationships.

Humans have the ability to comprehend, analyze, improve, and learn from the processes that occur around them. Animals, too, continue to improve the processes in which they are involved, whether consciously or unconsciously. Monkeys are known to create tools. The crow has been taught to collect trash in exchange for food. Over time, moths have developed techniques to avoid bats. During the drought season, an elephant herd can find water. Even trees in the tropics grow tall in order to compete for sunlight. There are numerous examples.

The improvement of processes in living creatures over time is a result of their behavior and evolution, rather than their intelligence. As time passes, methods are refined to produce better techniques and outcomes. Scholars and researchers have examined these tendencies toward process optimization and developed evolutionary optimization algorithms (EOAs) based on them. DMOA is one such EOA, where plants compete for sunlight and mycorrhizae compete for resources such as water, carbon, and zinc. The behavior patterns of these organisms are modeled in this algorithm [12].

The application of engineering practice, artificial intelligence, and managerial decision-making all include extensive research into optimization issues. The optimization problem gets more complicated as manufacturing size grows. The intricacy of the old exact algorithms based on derivative approaches limits their use in small-scale situations even though they can offer the exact answer to a problem [13].

Additionally, the problem that needs to be solved must have a model that is continuous and derivable; global optimization cannot be achieved for multi-peak, substantially nonlinear, or problems that are evolving dynamically [14]. As a result, solving global optimization problems via conventional techniques becomes very difficult.

Particle swarm optimization (PSO) [15], bat algorithm (BA) [16], JADE [17], cuckoo optimization algorithm (COA) [18], and flower pollination algorithm (FPA) [19] are examples of bioinspired intelligent optimization algorithms that have been proposed to solve complex optimization problems. The following algorithms simulate group behaviors and natural phenomena, such as the crow search algorithm (CSA) [20], grey wolf optimizer (GWO) [21], teaching–learning-based optimization (TLBO) [22], symbiotic organisms search (SOS) [23], and elephant herding optimization (EHO) [24]. Other cases are the YUKI algorithm, which is an innovative technique used to reduce the search space in specific problems. It involves creating a smaller local search area and simultaneously allocating two parts of the population for exploration and exploitation. This allows for the optimization of the search for interesting and promising solutions within the context of the problem being addressed [25–27], and the snake optimizer algorithm (SO) to address a diverse set of optimization tasks that mimics the special mating behavior of snakes. Each snake (male/female) struggles to have the best mate if the existing food is sufficient and the temperature is low, the algorithm mimics and mathematically models such behaviors and patterns of foraging and reproduction in a simple and efficient way [28]. These are only a few of the innovative intelligence algorithms that have been developed in the last five years. No intelligent algorithm, however, is capable of resolving every optimization issue. As a result, the freshly proposed or newly discovered algorithms have a wide range of application history.

Optimization is not restricted to a particular discipline such as applied mathematics, engineering, medicine, economics, computer science, or operations research. It has become a crucial tool in all areas of study. Advancements in developing novel algorithms and theoretical methods have enabled optimization to evolve in multiple directions, with a specific emphasis on artificial intelligence. This includes fields such as deep learning, machine learning, computer vision, fuzzy logic systems, and quantum computing [29,30].

Optimization has experienced a consistent expansion over the last five decades. Contemporary society not only exists in a fiercely competitive setting but also has to contemplate sustainable growth and conservation of resources. Therefore, it is crucial to optimally plan, design, operate, and manage resources and assets. Initially, the focus was on optimizing each operation independently. However, the present inclination is towards an integrated approach that encompasses synthesis and design, design and control, production planning, scheduling, and control [31].

Optimization has been evolving in recent years to the point that now provides general solutions to linear, nonlinear, unbounded, and constrained optimization problems. These problems are part of the mathematical programming area and can be divided into two classes: linear and nonlinear programming problems. Genetic algorithms and simulated annealing are two key methodologies that have been receiving increasing attention in real applications. The rapid development of technology has offered users a plethora of optimization codes with diverse degrees of rigor and complexity that can help in solving real-world problems. It is also possible to extend the capabilities of existing methods by integrating the features of two or more optimization methods to achieve a more efficient hybrid optimization algorithm [32]. However, there is no single method that can solve all specific problems of No Free Lunch (NFL) [33], and research is still ongoing to develop optimization methods that can solve them all.

Optimization methods have particular applications and are not applicable to every problem. It is necessary to recognize a problem as an optimization problem, or else other artificial intelligence techniques with specific specializations may be more appropriate.

3. Fuzzy Logic Systems

Fuzzy logic is a mathematical method that was introduced by Lotfi A. Zadeh in the 1960s [34] and is based on the theory of fuzzy sets. This theory proposes that an element can partially belong to a set, unlike in classical set theory where an element either belongs or does not belong to a set. This allows for an efficient way to work with uncertainties and to condition knowledge in the form of rules towards a quantitative level that can be processed by computers. Fuzzy logic is based on the way that people make decisions based on imprecise and linguistic information. Fuzzy sets are mathematical concepts for representing vagueness and imprecise information. The concept of fuzzy set membership is used to determine how much observation is within a set. Fuzzy logic deals with uncertainty in reasoning and utilizes concepts, principles, and methods developed within it.

Fuzzy logic is a computational technique that imitates the way humans think. Humans can make decisions based on vague information, such as determining if a room is hot without knowing the exact temperature. Fuzzy logic attempts to simulate this behavior of the human brain by using logical expressions that consider “degrees of truth” instead of the classic terms “true” or “false”. Equation (1) shows the equation for type-1 fuzzy logic systems (T1FLS), and Figure 2 shows the general scheme for T1FLS.

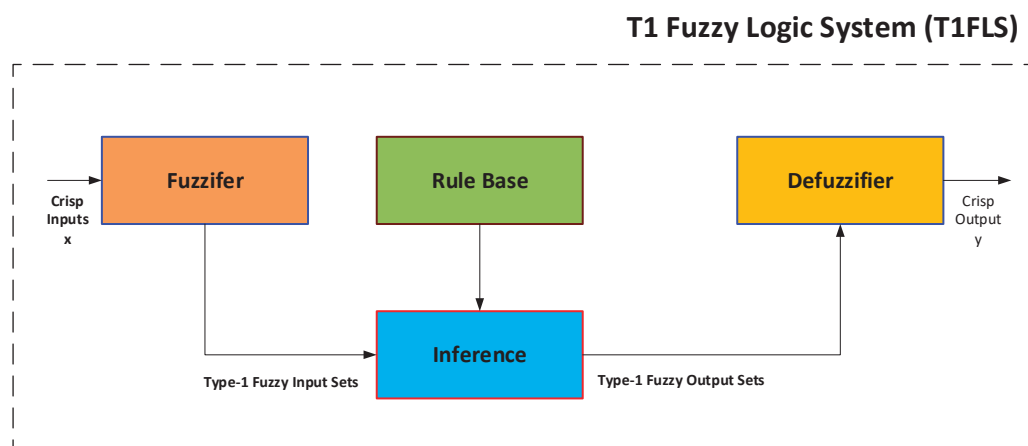


Figure 2. Scheme of a T1FLS.

The theory of fuzzy sets, developed by L. Zadeh, allows modeling the uncertainty that occurs in biological and social systems. Fuzzy logic is a theory dealing with uncertainty in reasoning and utilizes concepts, principles, and methods developed within it [35,36].

$$A = \{(x, \mu_A(x)) | x \in X\} \quad (1)$$

Interval Type-2 Fuzzy Logic System

Zadeh introduced the concept of type-2 fuzzy sets as an extension of fuzzy sets, specifically type-1 fuzzy sets, according to [37]. Type-2 fuzzy sets are characterized by having fuzzy membership degrees [38], which can be any subset of $[0, 1]$ of the primary membership. Additionally, for each primary membership, there is a corresponding secondary membership that defines the possibilities of the primary membership, and this secondary membership can range “between 0 and 1” as noted in [39]. Type-1 fuzzy sets are a special case of type-2 fuzzy sets, where the secondary membership function consists of a single element, the unit. The use of type-2 fuzzy sets enables us to manage linguistic uncertainty, such as the fact that “words can mean different things to different people”. Higher types of fuzzy relations, including type-2, increase fuzziness in relationships and can lead to greater logical processing of imprecise information, as stated in [40].

Interval type-2 fuzzy sets (IT2FLS) are a mathematical framework that builds on the original concepts of fuzzy sets, providing a means to account for uncertainty in models. Recent years have seen significant progress in this area. An IT2FLS can be defined mathe-

matically, as shown in [41,42]. A fuzzy set is a way to express the degree of truth for an element belonging to a set in a non-deterministic manner that allows for imprecision and uncertainty. In this context, a type-2 fuzzy set, denoted by $\tilde{A}$, is characterized by a type-2 membership function (x, u) where $x \in X$, $u \in J_x \subseteq [0, 1]$, and $0 \leq (x, u) \leq 1$ defined in Equation (2) and Figure 3.

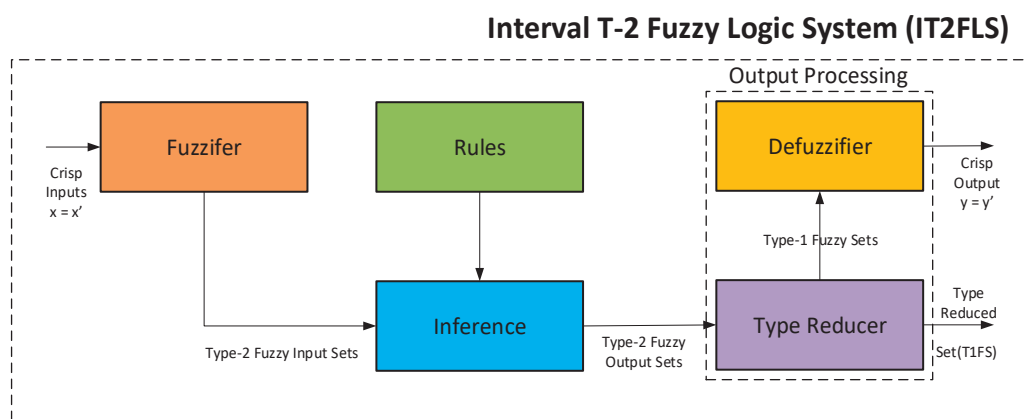


Figure 3. Scheme of an IT2FLS.

Fuzziness (entropy) is usually considered in measuring uncertainty for type-1 fuzzy sets, while for IT2FLS, centroid, and other measures are used to measure uncertainty.

$$\tilde{A} = \{((x, u), 1) | \forall x \in X, \forall u \in J_x \subseteq [0, 1]\} \quad (2)$$

4. Related Work

Nature-inspired optimization algorithms, such as particle Swarm Optimization (PSO) [43], flower pollination algorithm (FPA) [19,44], ant colony optimization (ACO) [45], artificial bee colony (ABC) [46], firefly algorithm (FA) [47], etc., have demonstrated flexibility, efficiency, and adaptability, in solving a wide spectrum of problems in real world applications, their merits and successes have inspired researchers to continuously develop these algorithms innovative.

For these reasons, it is proposed to develop an algorithm inspired by plants and how they adapt to physiological changes, survival, and growth through communication and the exchange of resources that are transferred through a fungal network. The proposed metaheuristics have a stochastic basis, that is, probabilistic, and the randomness rules are combined to imitate the process that inspires the algorithm.

Plant-inspired algorithms exist in the literature, it has been shown that plants exhibit intelligent behaviors “Plant intelligence-based metaheuristic optimization algorithms” [48], such as the one based on plant defense mechanism “A New Bio-inspired Optimization Algorithm Based on the Self-defense Mechanisms of Plants” [49], Flower Pollination Algorithm (FPA) “Fuzzy Flower Pollination Algorithm to Solve Control Problems” [50], Plant Growth Optimization (PGO) “A Global Optimization Algorithm Based on Plant Growth Theory: Plant Growth Optimization” [51].

Plants are highly successful in colonizing many habitats and represent approximately 99% of the planet’s eukaryotic biomass. They have evolved a variety of mechanisms to solve problems such as foraging and reproductive strategies. Plants sense environmental conditions and take measures to adapt to changing environments, such as searching for light and nutrients, to defend themselves against herbivores and other attackers. Although plants do not have a brain or central nervous system, they can sense environmental conditions and take “adaptive” measures that allow them to adjust to environmental changes. Plant adaptations are special features that improve their chances of survival and evolve over a long period of time. Examples of plant adaptations include:

- Foraging for light, water, and other nutrients;
- The ability to defend themselves against herbivores and other attackers;
- The ability to “remember” past events.

Plants and algae use photosynthesis to convert carbon dioxide and water into organic compounds, especially carbohydrates, using energy from sunlight and releasing oxygen as waste. Although plants are not known for their ability to move, they can move in response to various stimuli, but much more slowly than animals. Plants and algae are photosynthetic organisms that account for almost 50% of the photosynthesis that occurs on Earth. Photosynthesis is the process by which light energy is converted to chemical energy, whereby carbon dioxide and water are converted into organic molecules. Plants can move in response to a variety of stimuli, which include:

1. Light (phototropism), plants constantly monitor their visible environment.
2. Gravity (geotropism), the plant’s root network also moves, and the root tips respond to gravity.
3. Water (hydrotropism), which is the response of plant growth to water.
4. Touch (thigmotropism), many plants respond to the sense of touch, such as the tendrils of climbing plants, vines, or bindweed.

With plant propagation algorithms, plants have a propagation process, such as seed dispersal and root propagation. The invasive weed optimization algorithm (IWO), based on the colonization behavior of weeds, was put forward by Mehrabian and Lucas (2006) [52]. The paddy field algorithm was first proposed by Premaratne, Samarabandu, and Sidhu (2009) [53], and is inspired by aspects of the plant reproduction cycle, focusing on pollination and seed dispersal process. Although many plants are propagated using seeds, some employ a system of “runners” or horizontal stems that grow outward from the base of the plant. The strawberry plant algorithm is inspired by the propagation of plants through seeds and stolons [54]. In the plant growth simulation algorithm (PGSA), inspired by the light foraging process, an important aspect of plant growth is that the initial plant stem eventually gives rise to branches and leaves as it is growing [55].

Despite the wide variety of plants and associated plant behaviors that occur in the natural world, little inspiration has so far been taken from these mechanisms for the design of computational algorithms.

So far, there is no algorithm with the characteristics that mimic the behavior of an ecosystem such as a forest and specifically in the understory, i.e., the behavior between tree roots and a fungal network.

5. Proposed Method

The novel DMOA algorithm is inspired by the nature of the Mycorrhiza Network (MN) and plant roots with this close interaction between these two organisms (plant roots and MN fungal network), a symbiosis is generated, and it has been found that in this relationship [56–60]:

- There is communication among plants, which may or may not be of the same species, through a fungal network (MN).
- There is an exchange of resources among plants through the fungal network (MN).
- There is a defense behavior against predators that can be insects or animals, for the survival of the whole habitat (plants and fungi).
- The colonization of a forest through a fungal network (MN) thrives much more than in a forest where there is no exchange of resources (see Figure 4).

This new optimization method FDMOA inspired by the symbiosis of plants and the Mycorrhizal Network uses the six discrete Lotka–Volterra system equations (DLVSE), these equations model the understory ecosystem where plant roots and the MN have a symbiotic relationship. With Equations (3) and (4) in predator–predator model, biochemical signals that travel through the MN alert all the plants that are connected to this network to the danger of predators, fires, floods, etc., with Equations (5) and (6) in the cooperative model,

it transfers resources from plants to other growing plants and from plants to the MN, all these resources travel through the MN and Equations (7) and (8) competitive model, it competes for habitat resources with respect to other plants for obtaining sunlight to perform photosynthesis that is converted into carbon that they share with the MN, the water and minerals that the MN obtains are shared with the plants.

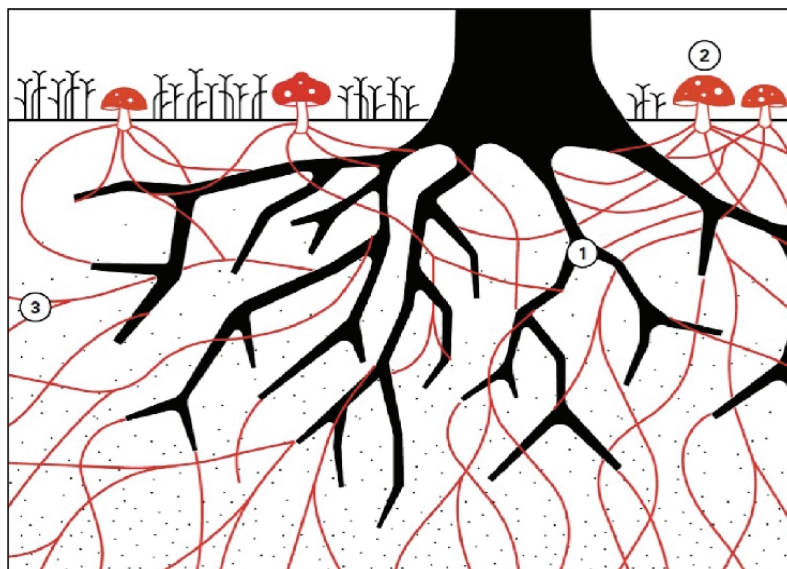


Figure 4. Nutrient transport through the MN.

It has been demonstrated that the algorithm has a fast convergence and therefore a low computational cost, the three biological operators represented by DLVSE are the defense model, the cooperative model and the competitive model, initially we have two populations (plants and MN), both populations are obtained by generating random numbers, we obtain the best fitness of each population, these values are the input to the parameters a (population grow rate x) and d (population grow rate y) of the DLVSE equation system, the result of Equation (9) (iterations), is the input for the fuzzy systems T1FLS and IT2FLS, the parameterization of the membership functions are modified with values provided by the FDMOA algorithm, the output of the fuzzy system is the parameter ξ (growth rates of the populations x in time t) that influence in a determinant way the convergence of the algorithm, then the biological operator resource exchange (cooperative model) and its result has inference in one of the two biological operators (defensive model or competitive model) based on a random outcome 1 or 2, with this we try to simulate what happens in an ecosystem where these events of defense against predators and competition for resources such as water, carbon, zinc, etc. Then the fitness is evaluated, the population and fitness are updated, and the stop condition is verified, if it is higher the algorithm is terminated, otherwise the process continues with the fuzzy systems process. The sequence of the FDMOA algorithm can be seen in Algorithm 1 showing the pseudocode and Figure 5 illustrating the flowchart of the FDMOA algorithm.

All the parameters in Algorithm 1, have been obtained from the literature [61–64] and are sensitive to the results obtained in this investigation, the only parameters that we experiment that move in an important way the generation of new values are the parameters a and d that we mentioned previously, the parameter ξ generated by the fuzzy systems, is the parameter that moves the fuzzy system towards the convergence, we have to investigate in depth each one of the parameters to see its incidence in the results; this investigation will be reason to make another article.

Algorithm 1: Fuzzy Discrete Mycorrhiza Optimization Algorithm (FDMOA)

```

1:  Objective min or max  $f(x)$ ,  $x = (x_1, x_2, \dots, x_d)$ 
2:  Define parameters  $(a, b, c, d, e, f, x, y)$ 
3:  Initialize a population of  $n$  plants and mycorrhiza with random solutions
4:  Find the best solution fit in the initial population
5:  while ( $t < \text{maxIter}$ )
6:    for  $i = 1:n$  (for  $n$  plants and Mycorrhiza population)
7:       $X_p = \text{abs}(\text{FitA})$ 
8:       $X_m = \text{abs}(\text{FitB})$ 
9:    end for
10:    $a = \text{minor}X_p$ 
11:    $d = \text{minor}X_m$ 
12:   Apply (LV-Cooperative Model)
13:    $x_i^{t+1} = \frac{(ax_i - bx_i y_i)}{(1 - gx_i)}$ 
14:    $y_i^{t+1} = \frac{(dy_i + ex_i y_i)}{(1 + hy_i)}$ 
15:   if  $x_i < y_i$ 
16:      $x^t = x_i$ 
17:   else
18:      $x^t = y_i$ 
19:   end if
20:    $\text{rand}([1 \ 2])$ 
21:   if ( $\text{rand} = 1$ )
22:     Apply (LV-Predator-Prey Model)
23:      $x_i^{t+1} = ax_i(1 - x_i) - bx_i y_i$ 
24:      $y_i^{t+1} = dx_i y_i - gy_i$ 
25:   else
26:     Apply (LV-Competitive Model)
27:      $x_i^{t+1} = \frac{(ax_i - bx_i y_i)}{(1 + gx_i)}$ 
28:      $y_i^{t+1} = \frac{(dy_i - ex_i y_i)}{(1 + hy_i)}$ 
29:   end if
30:   Evaluate new solutions.
31:   T1FLS-IT2FLS Architecture
32:   Evaluate Error
33:   Error minor?
34:   Update T1FLS-IT2FLS Architecture.
35:   Find the current best FLS-Architecture solution.
36: end while

```

By dynamic adaptation of parameters in this method we refer to the change in the parameter values of the membership functions of the T1FLS and IT2FLS fuzzy systems in each iteration with the purpose of improving the performance and precision of the algorithm.

This new optimization method inspired by the symbiosis of plants and the Mycorrhiza Network, the FDMOA algorithm uses the discrete Lotka–Volterra system equations (DLVSE), it has been demonstrated that the algorithm has fast convergence and therefore a low computational cost, the three biological operators represented by DLVSE are the defense model, cooperative model, and competitive model, initially we have two populations (plants and MN), we obtain the best fitness of each population, the result of Equation (9) (iterations), is the input for the fuzzy systems T1FLS and IT2FLS, the parameterization of the membership functions are modified with values provided by the FDMOA algorithm, the output of the fuzzy system is the parameter x_i (grow rates of populations x at time t) that influence of determinant form in the convergence of the algorithm, then the biological operator resource exchange (cooperative model) and its result has inference in one of the two biological operators (defense model or competitive model) based on a random result 1 or 2, with this we try to simulate what happens in an ecosystem where these events of defense against predators and competition for resources such as water, carbon, zinc, etc.,

then the fitness is evaluated, the population and fitness are updated, the stop condition is checked, if it is higher the algorithm is terminated, otherwise the process continues in the fuzzy systems step. The sequence of the FDMOA algorithm can be seen in the pseudocode of Algorithm 1 and in the flowchart of Figure 5.

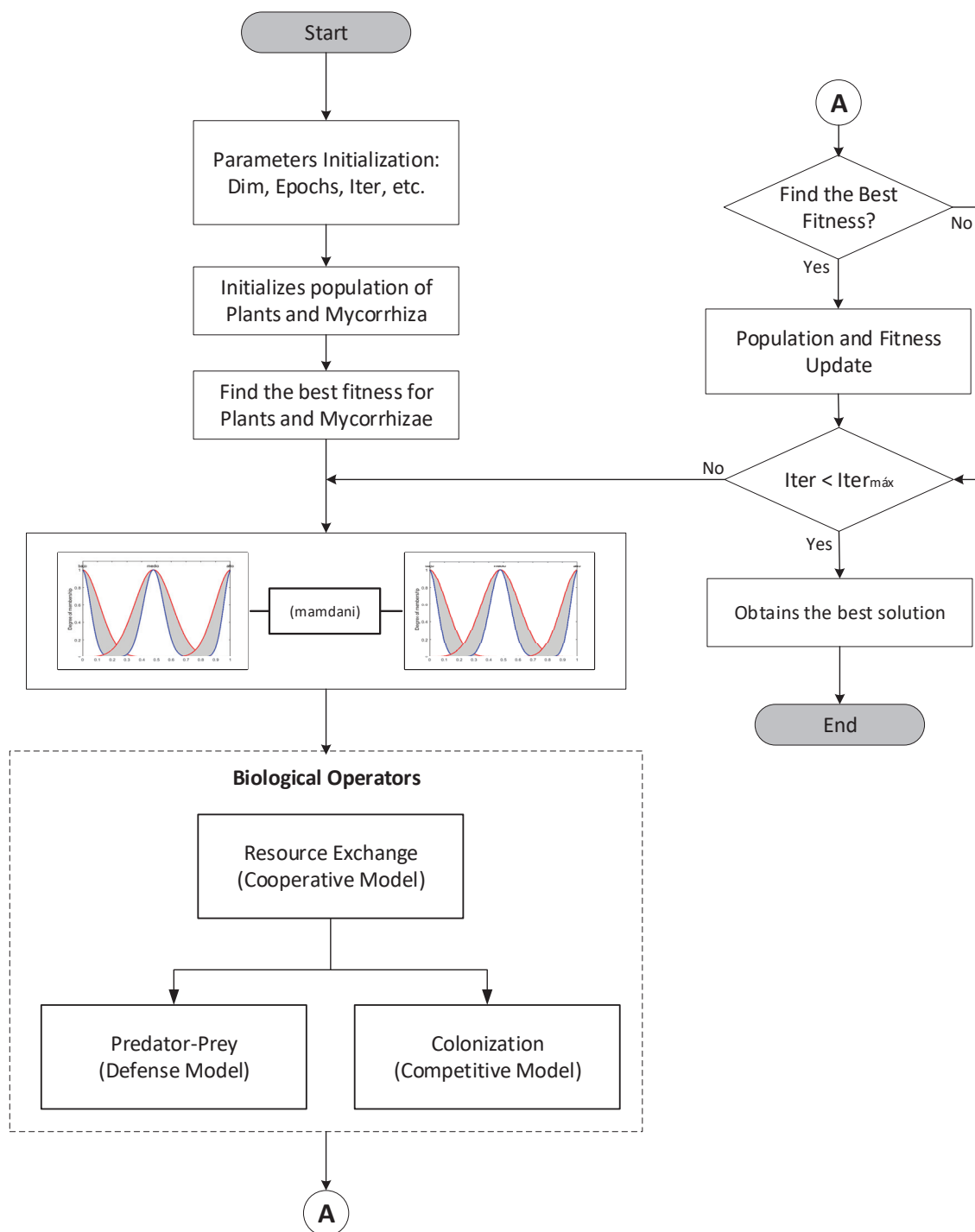


Figure 5. FDMOA flowchart.

5.1. Discrete Mycorrhiza Optimization Algorithm

The mycorrhizal associations between plants and fungi have significant impacts on the ecosystem at a large scale. This is mostly due to the fact that most plants tend to form these associations, which are believed to have originated in ancient times and helped plants

colonize the land. The symbiosis between plants and fungi is a many-to-many relationship, meaning that plants can form associations with a wide variety of fungal species, and fungal species can colonize many different plant species. While most mycorrhizal fungi have a broad range of hosts and form diffuse mutualisms, some are specialists that only occur in one host.

It is now understood that the Mycorrhiza Network (MN) can impact various aspects of plant life, including establishment, survival, physiology, growth, and chemical defense. This impact is thought to occur because MN serves as a pathway for exchanging stress molecules and resources between plants. For instance, the most common method for mycorrhizal fungal colonization of regenerating plants in their natural environment is believed to be through anastomosis with pre-existing MN of established plants. The colonization of seedlings by MN enables them to obtain enough nutrients from the soil for the growth of their roots and shoots, leading to their survival.

The behavior of living things in nature has inspired researchers in computer science to develop new optimization algorithms, with a focus on the relationship that mycorrhiza fungi have developed with the roots of plants, specifically trees. The colonization of trees on the earth would not have occurred except for the mycorrhiza fungal networks through their roots. To date, 100,000 species of fungi are known, but it is possible that there are more. This relationship between fungi and plants is an example of symbiosis, where there is a mutual exchange of resources between the two organisms, with fungi providing plants with nitrates and phosphates necessary for their growth in exchange for carbon dioxide carried out through photosynthesis, resulting in mutual benefit and improved biological fitness for both plant and fungus [65].

Symbiosis refers to a close and ongoing biological interaction between different species of organisms. In the case of fungi, they reside either on the surface of the roots or within the bark of plant roots, as illustrated in Figure 4. This interaction involves an exchange of resources between the fungi and plants, where the fungi provide nitrates and phosphates that are essential for plant growth in return for carbon dioxide produced through photosynthesis. This leads to a mutually beneficial relationship, or mutualism, where both the plant and fungus benefit and enhance their biological fitness [66–68].

In Figure 4 we can see that a mycorrhiza is a symbiotic relationship between roots (1) and fungi (2) that involves the exchange of plant and tree sugars for moisture and nutrients acquired by fungal filaments (3) from the soil. By extending the root systems of trees, mycorrhizae significantly improve their absorptive capacity, expanding their ability to gather essential resources.

5.2. Discrete Lotka–Volterra System Equation

The discrete Lotka–Volterra system equations (DLVSE) used in this research are Linear Equations (3)–(8) which are described below:

Equations (3) and (4) [61,62], were used to develop the biological operator (predator–prey model) within the algorithm.

Equations (5) and (6) [63], were used to develop the biological operator (cooperative model) in the algorithm.

Equations (7) and (8) [63,64], were used to develop the biological operator (competitive model) in the algorithm.

$$x_i^{t+1} = \frac{(ax_i - bx_i y_i)}{(1 - gx_i)} \quad (3)$$

$$y_i^{t+1} = \frac{(dy_i + ex_i y_i)}{(1 + hy_i)} \quad (4)$$

$$x_i^{t+1} = ax_i(1 - x_i) - bx_i y_i \quad (5)$$

$$y_i^{t+1} = dx_i y_i - gy_i \quad (6)$$

$$x_i^{t+1} = \frac{(ax_i - bx_i y_i)}{(1 + gx_i)} \quad (7)$$

$$y_i^{t+1} = \frac{(dy_i - ex_i y_i)}{(1 + hy_i)} \quad (8)$$

5.3. FDMOA Parameters

Table 1 provides all the parameters that are used in the FDMOA algorithm, such as populations, dimensions, epochs, iterations, etc. These parameters are fixed, but we also consider some parameters as dynamic, as is explained later.

Table 1. FDMOA parameters.

Parameter	Description	Value
DMOA—Parameters:		
x_i^{t+1}	Population x at time t	
y_i^{t+1}	Population y at time t	
x_i	Grow rates of populations x at time t	
y_i	Grow rates of populations y at time t	
t	time	
a	Population growth rate x	0.01
b	Influence of population x on itself	0.02
g	Influence of population y on population x	0.06
d	Population growth rate y	0
e	Influence of population x on population y	1.7
h	Influence of population y on itself	0.09
x	Initial population in x	0.0002
y	Initial population in y	0.0006
In the absence of population $x = 0$, In the absence of population $y = 0$		
a, b, c, d, e and f —are positive constants		
Population	Population size	20
Populations	Number of populations	2
Dimensions	Dimensions size	30, 50, 100
Epochs	Number of epochs	30
Iterations	Iteration's size	30, 50, 100, 500

5.4. FDMOA Pseudocode

Algorithm 1 shows the pseudocode of the logic and structure of the FDMOA optimization algorithm.

5.5. FDMOA Flowchart

Figure 5 shows the process flow diagram of the FDMOA optimization algorithm.

5.6. Mathematical Functions

Table 2 contains the 36 mathematical functions with which we performed all the experimentation in this article: function number (F), function name, range, and nature (U: unimodal or M: multimodal).

Figures 6–9 represent the graphical schemes of the 36 different mathematical functions (CEC2013) and Figure 10 shows the 36 mathematical functions, mentioned in Table 2, which we are using for the experimentation of the DMOA algorithm, in the figures we can find the function number (F), the image of the function.

Sphere function is a convex and unimodal function, which means that it has a single global minimum, optimization algorithms usually have no problem finding the global minimum. Rosenbrock function is a non-convex function with a long narrow valley leading to the global minimum. Griewank function is a multimodal function with multiple local minima, it presents a challenging optimization landscape with oscillations. Rastrigin function is a highly multimodal function with a rough landscape, its optimization challenges

arise from the presence of many local minima. Ackley function is a multimodal function with a complex landscape, its optimization challenges stem from the trade-off between exploration and exploitation. Dixon-Price function is a unimodal function with multiple local minima. Michalewicz function is a multimodal function with a complex landscape, it tests the ability of an optimization algorithm to handle high-dimensional problems with intricate relationships between variables. Powell function is a mathematical function with multiple local minima, its optimization challenges stem from its high dimensionality and the presence of local minima. The Rotate Hyper-Ellipsoid function is a smooth elongated function used to evaluate optimization algorithms. Schwefel function has a very rough landscape with multiple local minima, Figure 6.

Table 2. 36 Mathematical functions.

F	Function	Range	Nature
F1	Sphere	$[-5.12, 5.12]$	U
F2	Rosenbrock	$[-5, 10]$	U
F3	Griewank	$[-600, 600]$	M
F4	Rastrigin	$[-5.12, 5.12]$	M
F5	Ackley	$[-32.768, 32.768]$	M
F6	Dixon-Price	$[-10, 10]$	U
F7	Michalewicz	$[0, \pi]$	M
F8	Powell	$[-4, 5]$	U
F9	RHE: Rotate Hyper Ellipsoid	$[-65.536, 65.536]$	U
F10	Shwefel	$[-500, 500]$	M
F11	Styblinski–Tang	$[-5, 5]$	U
F12	SDP: Sum Different Powers	$[-1, 1]$	M
F13	Sum Squares	$[-10, 10]$	U
F14	Trid	$[-d^2, d^2]$	U
F15	Zakharov	$[-5, 10]$	U
F16	Bukin No 6	$[-15, -5]$	U
F17	Cross-in-Tray	$[-10, 10]$	M
F18	Drop-Wave	$[-5.12, 5.12]$	M
F19	Eggholder	$[-5.12, 5.12]$	M
F20	Beale	$[-4.5, 4.5]$	U
F21	Holder Table	$[-10, 10]$	M
F22	Branin	$[-5, 10]$	M
F23	Levy	$[-10, 10]$	M
F24	Levy 13	$[-10, 10]$	M
F25	Schaffer 2	$[-100, 100]$	M
F26	Schaffer 4	$[-100, 100]$	M
F27	Shubert	$[-10, 10]$	M
F28	Bohachevsky 1	$[-100, 100]$	M
F29	Bohachevsky 2	$[-100, 100]$	M
F30	Bohachevsky 3	$[-100, 100]$	M
F31	Booth	$[-10, 10]$	U
F32	Matyas	$[-10, 10]$	U
F33	Mccormick	$[-1.5, 4]$	U
F34	Easom	$[-100, 100]$	U
F35	Goldstein–Price	$[-2, 2]$	M
F36	Three-Hump Camel	$[-5, 5]$	M
U	Unimodal		
M	Multimodal		

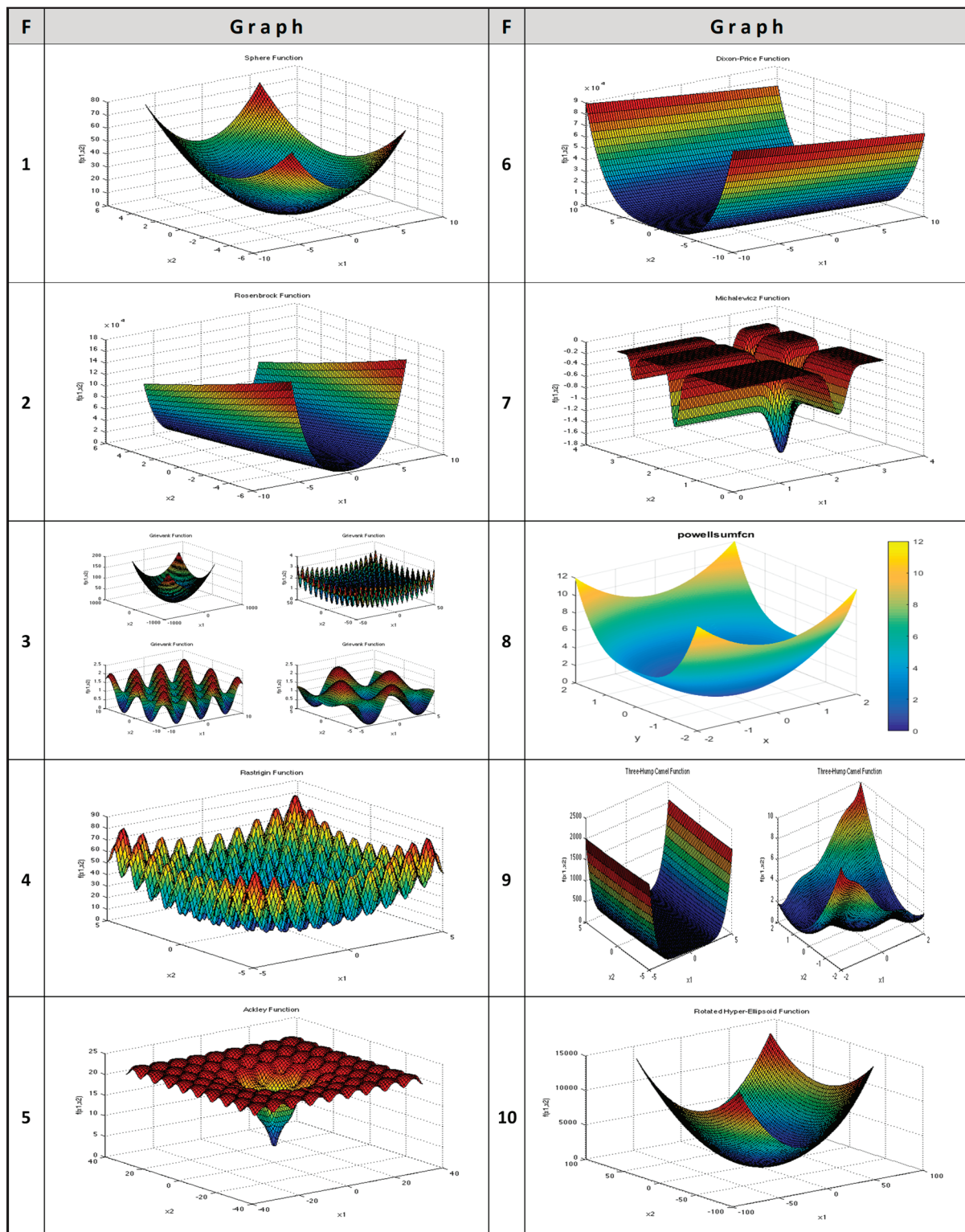


Figure 6. Graphs of the mathematical functions: Sphere, Rosenbrock, Griewank, Rastrigin, Ackley, Dixon, Michalewicz, Powell, Rotate Hyper Ellipsoid, and Schwefel.

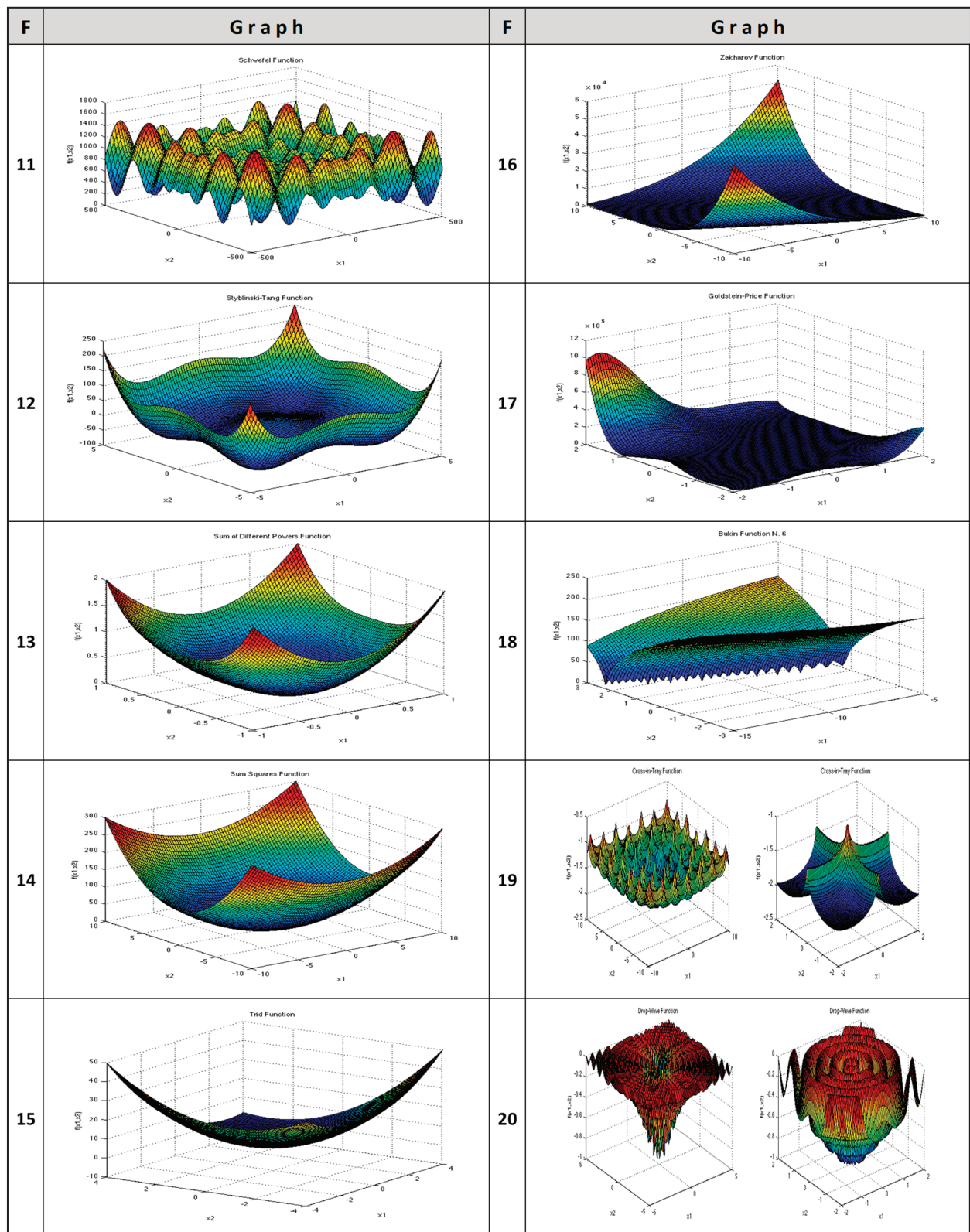


Figure 7. Graphs of the mathematical functions: Styblinski–Tang, Sum Different Powers, Sum Squares, Trid, Zakharov, Bukin No 6, Cross-in-Tray, Drop-Wave, Eggholder, and Beale.

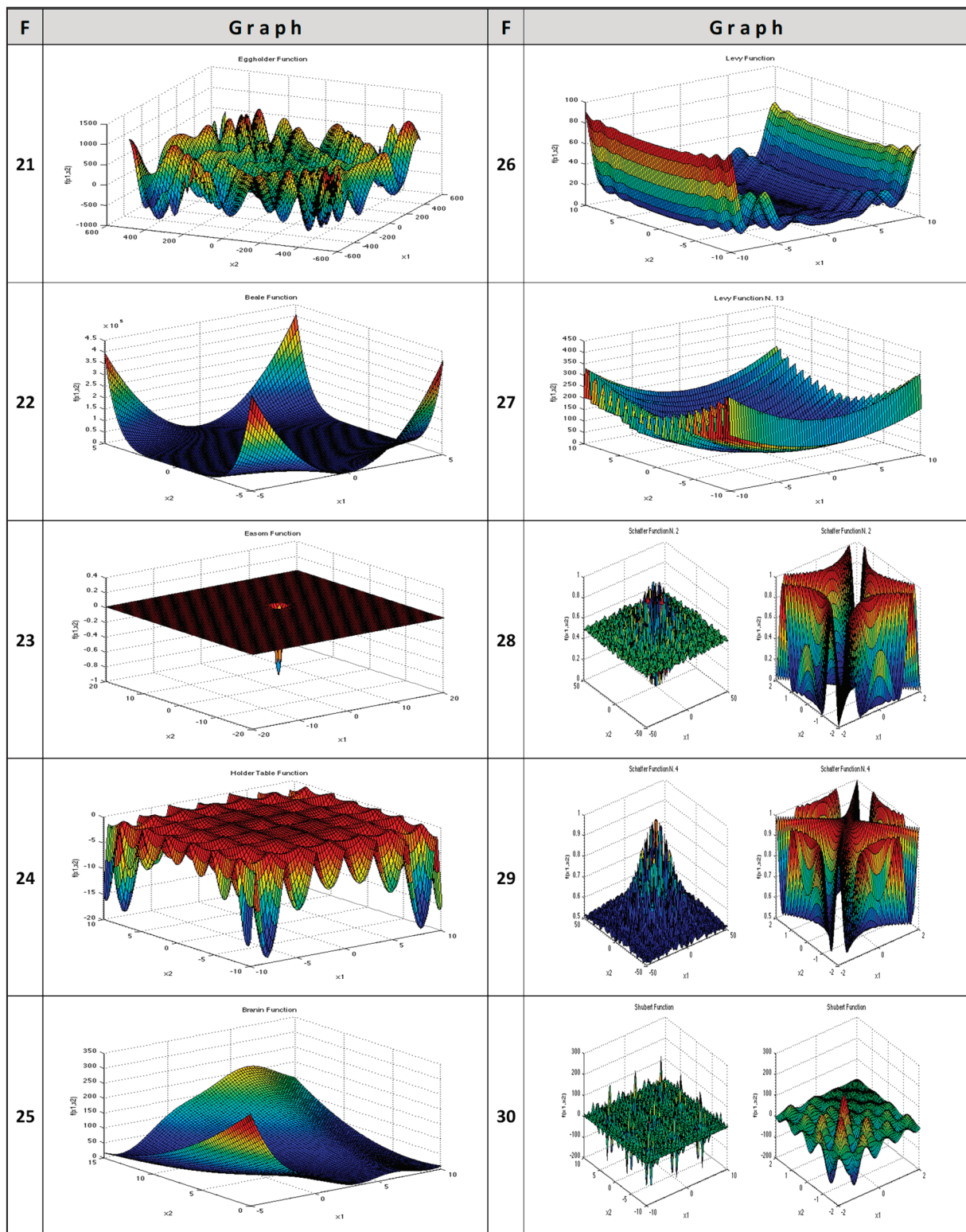


Figure 8. Graphs of the mathematical functions: Holder Table, Branin, Levy, Levy 13, Shaffer 2, Shaffer 4, Shubert, Bohachevsky 1, Bohachevsky 2, and Bohachevsky 3.

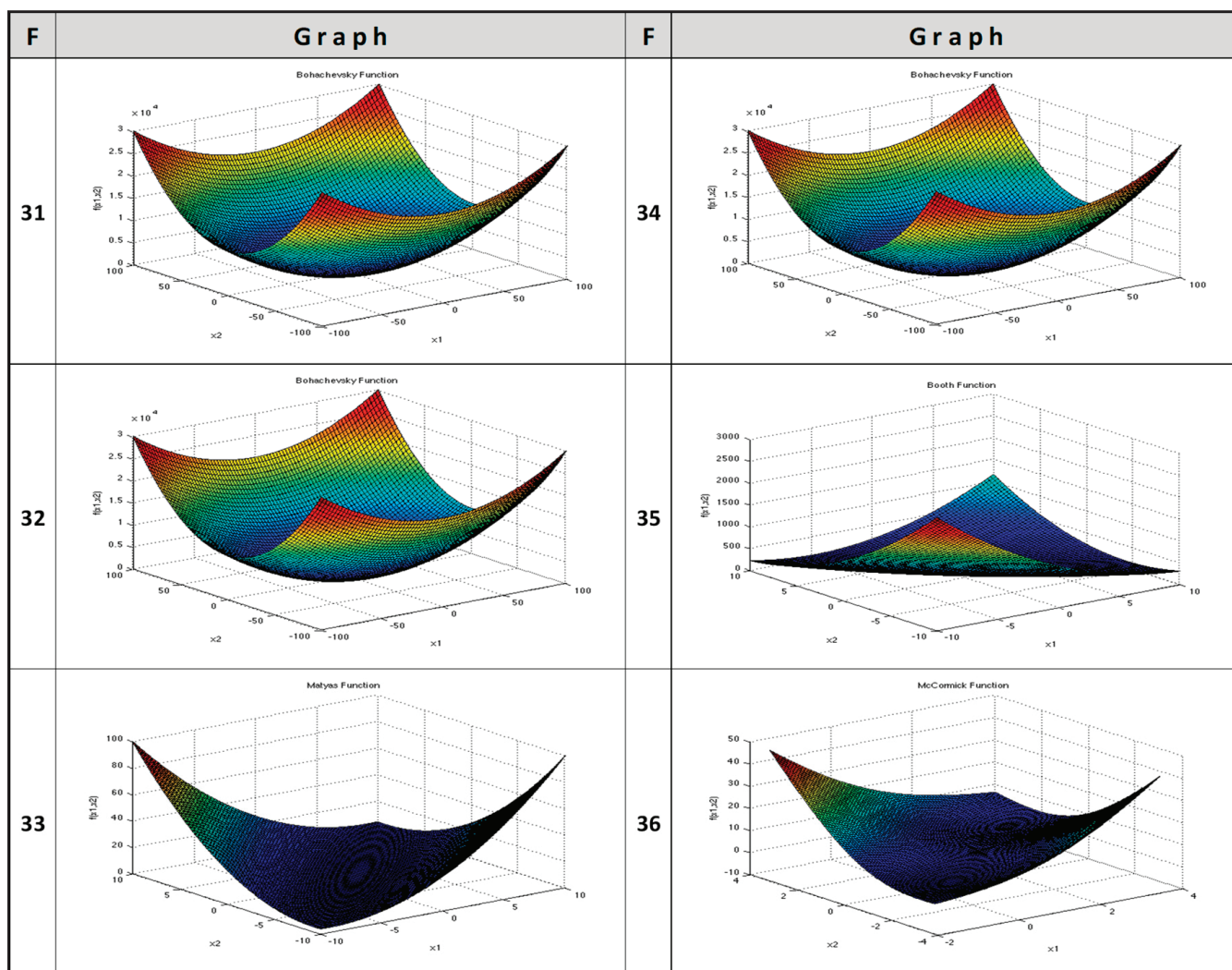


Figure 9. Graphs of the mathematical functions: Booth, Matyas, McCormick, Easom, Goldstein–Price, and Three-Hump Camel.

Styblinski-Tang function is a unimodal function with many local minima. Sum Different Powers function is a multimodal function with multiple local minima, it presents challenges in optimization due to its hilly landscape and the presence of local minima. Sum Squares function is a convex and unimodal function, optimization algorithms usually have no problems finding the global minimum. Trid function is a unimodal function with multiple local minima. Zakharov function is a unimodal function with multiple local minima, its optimization problems arise from its hilly landscape and the presence of local minima. Bukin No 6 function is a unimodal function with multiple local minima, its optimization problems arise from the complexity and irregularity of the environment. Cross-in-Tray function is a multimodal function with multiple local minima, its optimization problems lie in the presence of many local minima and in the complexity of the environment. Drop-Wave function is a multimodal function with a challenging optimization landscape, it has a global minimum in a narrow region surrounded by many local minima. Eggholder function is a multimodal function with a complex landscape, its optimization problems arise from the presence of multiple local minima and intricate relationships between variables. Beale function is a unimodal function with multiple local minima, optimization algorithms face the challenge of efficiently exploring the search space, Figure 7.

N	Function	N	Function
1	$f(\mathbf{x}) = \sum_{i=1}^d x_i^2$	11	$f(\mathbf{x}) = (x_1 - 1)^2 + \sum_{i=2}^d i (2x_i^2 - x_{i-1})^2$
2	$f(\mathbf{x}) = \sum_{i=1}^{d-1} [100(x_{i+1} - x_i^2)^2 + (x_i - 1)^2]$	12	$f(\mathbf{x}) = -\sum_{i=1}^d \sin(x_i) \sin^{2m}\left(\frac{ix_i^2}{\pi}\right)$
3	$f(\mathbf{x}) = \sum_{i=1}^d \frac{x_i^2}{4000} - \prod_{i=1}^d \cos\left(\frac{x_i}{\sqrt{i}}\right) + 1$	13	$f(\mathbf{x}) = \sum_{i=1}^{d/4} [(x_{4i-3} + 10x_{4i-2})^2 + 5(x_{4i-1} - x_{4i})^2 + (x_{4i-2} - 2x_{4i-1})^4 + 10(x_{4i-3} - x_{4i})^4]$
4	$f(\mathbf{x}) = 10d + \sum_{i=1}^d [x_i^2 - 10 \cos(2\pi x_i)]$	14	$f(\mathbf{x}) = 2x_1^2 - 1.05x_1^4 + \frac{x_1^6}{6} + x_1x_2 + x_2^2$
5	$f(\mathbf{x}) = -a \exp\left(-b \sqrt{\frac{1}{d} \sum_{i=1}^d x_i^2}\right) - \exp\left(\frac{1}{d} \sum_{i=1}^d \cos(cx_i)\right) + a + \exp(1)$	15	$f(\mathbf{x}) = \sum_{i=1}^d \sum_{j=1}^i x_j^2$
6	$f(\mathbf{x}) = 418.9829d - \sum_{i=1}^d x_i \sin(\sqrt{ x_i })$	16	$f(\mathbf{x}) = \sum_{i=1}^d x_i^2 + \left(\sum_{i=1}^d 0.5ix_i\right)^2 + \left(\sum_{i=1}^d 0.5ix_i\right)^4$
7	$f(\mathbf{x}) = \frac{1}{2} \sum_{i=1}^d (x_i^4 - 16x_i^2 + 5x_i)$	17	$f(\mathbf{x}) = [1 + (x_1 + x_2 + 1)^2(19 - 14x_1 + 3x_1^2 - 14x_2 + 6x_1x_2 + 3x_2^2)] \times [30 + (2x_1 - 3x_2)^2(18 - 32x_1 + 12x_1^2 + 48x_2 - 36x_1x_2 + 27x_2^2)]$
8	$f(\mathbf{x}) = \sum_{i=1}^d x_i ^{i+1}$	18	$f(\mathbf{x}) = 100\sqrt{ x_2 - 0.01x_1^2 } + 0.01 x_1 + 10 $
9	$f(\mathbf{x}) = \sum_{i=1}^d ix_i^2$	19	$f(\mathbf{x}) = -0.0001 \left(\sin(x_1) \sin(x_2) \exp\left(\left 100 - \frac{\sqrt{x_1^2 + x_2^2}}{\pi}\right \right) + 1 \right)^{0.1}$
10	$f(\mathbf{x}) = \sum_{i=1}^d (x_i - 1)^2 - \sum_{i=2}^d x_i x_{i-1}$	20	$f(\mathbf{x}) = -\frac{1 + \cos\left(12\sqrt{x_1^2 + x_2^2}\right)}{0.5(x_1^2 + x_2^2) + 2}$
21	$f(\mathbf{x}) = -(x_2 + 47) \sin\left(\sqrt{\left x_2 + \frac{x_1}{2} + 47\right }\right) - x_1 \sin\left(\sqrt{ x_1 - (x_2 + 47) }\right)$	26	$f(\mathbf{x}) = \sin^2(\pi w_1) + \sum_{i=1}^{d-1} (w_i - 1)^2 [1 + 10 \sin^2(\pi w_i + 1)] + (w_d - 1)^2 [1 + \sin^2(2\pi w_d)]$, where $w_i = 1 + \frac{x_i - 1}{4}$, for all $i = 1, \dots, d$
22	$f(\mathbf{x}) = (1.5 - x_1 + x_1x_2)^2 + (2.25 - x_1 + x_1x_2^2)^2 + (2.625 - x_1 + x_1x_2^3)^2$	27	$f(\mathbf{x}) = \sin^2(3\pi x_1) + (x_1 - 1)^2 [1 + \sin^2(3\pi x_2)] + (x_2 - 1)^2 [1 + \sin^2(2\pi x_2)]$
23	$f(\mathbf{x}) = -\cos(x_1) \cos(x_2) \exp(-(x_1 - \pi)^2 - (x_2 - \pi)^2)$	28	$f(\mathbf{x}) = 0.5 + \frac{\sin^2(x_1^2 - x_2^2) - 0.5}{[1 + 0.001(x_1^2 + x_2^2)]^2}$
24	$f(\mathbf{x}) = -\left \sin(x_1) \cos(x_2) \exp\left(\left 1 - \frac{\sqrt{x_1^2 + x_2^2}}{\pi}\right \right) \right $	29	$f(\mathbf{x}) = 0.5 + \frac{\cos(\sin(x_1^2 - x_2^2)) - 0.5}{[1 + 0.001(x_1^2 + x_2^2)]^2}$
25	$f(\mathbf{x}) = a(x_2 - bx_1^2 + cx_1 - r)^2 + s(1 - t)\cos(x_1) + s$	30	$f(\mathbf{x}) = \left(\sum_{i=1}^5 i \cos((i+1)x_1 + i)\right) \left(\sum_{i=1}^5 i \cos((i+1)x_2 + i)\right)$
31	$f_1(\mathbf{x}) = x_1^2 + 2x_2^2 - 0.3\cos(3\pi x_1) - 0.4\cos(4\pi x_2) + 0.7$	34	$f_2(\mathbf{x}) = x_1^2 + 2x_2^2 - 0.3\cos(3\pi x_1)\cos(4\pi x_2) + 0.3$
32	$f_3(\mathbf{x}) = x_1^2 + 2x_2^2 - 0.3\cos(3\pi x_1 + 4\pi x_2) + 0.3$	35	$f(\mathbf{x}) = (x_1 + 2x_2 - 7)^2 + (2x_1 + x_2 - 5)^2$
33	$f(\mathbf{x}) = 0.26(x_1^2 + x_2^2) - 0.48x_1x_2$	36	$f(\mathbf{x}) = \sin(x_1 + x_2) + (x_1 - x_2)^2 - 1.5x_1 + 2.5x_2 + 1$

Figure 10. 36 mathematical functions.

Holder Table function is a multimodal function with multiple local minima, it has a complex and irregular landscape with distinct peaks and valleys. Branin function is a multimodal function with multiple local minima. Levy function is a multimodal function

with a complex and challenging landscape. Levy 13 function is a multimodal function with multiple local minima, it presents optimization challenges due to its complex and irregular landscape. Shaffer 2 function is a multimodal function with multiple local minima, it has a complex and hilly landscape, which makes it challenging for optimization algorithms. Shaffer function 4 is a multimodal function with multiple local minima, it has a complex and irregular landscape. Shubert function is a multimodal function with a highly oscillatory landscape, it has multiple local minima and a global minimum. Bohachevsky 1 function is a multimodal function with multiple local minima, it has a complex and irregular landscape. Bohachevsky function 2 is a multimodal function with multiple local minima, it has a complex and irregular landscape. Bohachevsky function 3 is a multimodal function with multiple local minima, it presents optimization challenges due to its complex and irregular landscape, Figure 8.

Booth function is a unimodal function with multiple local minima, it has a simple but narrow valley landscape. Matyas function is a unimodal function with multiple local minima, it has a bowl-shaped landscape. McCormick function is a unimodal function with multiple local minima, it has a complex and irregular landscape. The Easom function is a unimodal function with multiple local minima, it has a sharp, narrow peak surrounded by a hilly landscape. The Goldstein-Price function is a multimodal function with multiple local minima, it has a complex and oscillatory landscape. Three-Hump Camel function is a multimodal function with multiple local minima, it has three distinct peaks and valleys, Figure 9.

Table 3 shows the three fuzzy IF THEN rules that were used for both the T1FLS and IT2FLS fuzzy systems.

Table 3. Rules IF THEN for T1FLS and IT2FLS.

N	Rules If Then
1	if (iter is Low) then (x_i is High)
2	if (iter is Medium) then (x_i is Medium)
3	if (iter is High) then (x_i is Low)

Equation (9) shows the way to calculate the “iteration” variable:

$$Iteration = \frac{Current\ Iteration}{Total\ Iterations} \quad (9)$$

Figure 11 shows the architecture for T1FLS Fis-fisGau318 with parameter adaptation, an input with three Gaussian functions and output also with three Gaussian functions using the Mamdani method, Gaussian membership function “Low” with blue color, Gaussian membership function “Medium” with orange color and Gaussian membership function “High” with yellow color, “iter” is the number of iteration parameters and x_i is a DLVSE parameter indicating grow rates of populations x at time t .

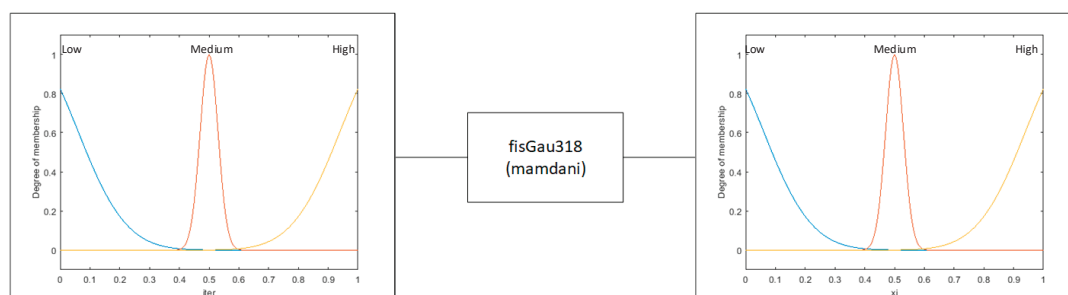


Figure 11. Architecture for T1FLS FIS-fisGau318.

Figure 12 shows the architecture for the IT2FLS FIS-it2_gausS01 with parameter adaptation, an input with three Gaussian functions, and an output also with three Gaussian type-2 functions using the Mamdani method, in the three Gaussian membership functions Low, Medium and High, the red line is the value of the upper membership function (UMF), the blue line is the value of the lower membership function (LMF) and the The internal part of the membership function with the gray color is called the footprint of uncertainty (FOU).

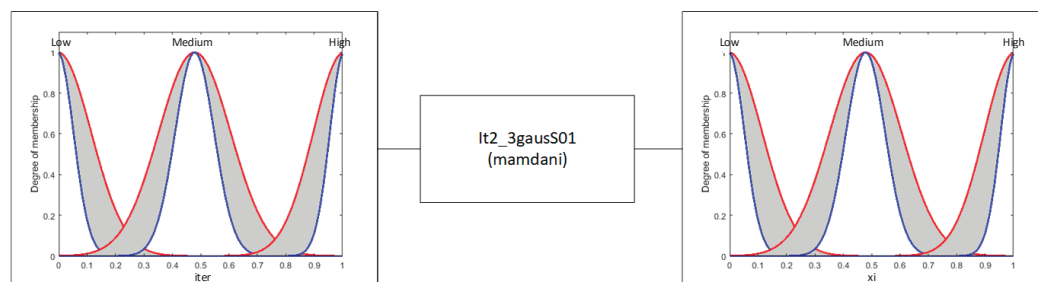


Figure 12. Architecture for IT2FLS FIS-it2_3gausS01.

Figure 13 shows the architecture for IT2FLS FIS-it2_3gausS6523 optimized with the DMOA algorithm.

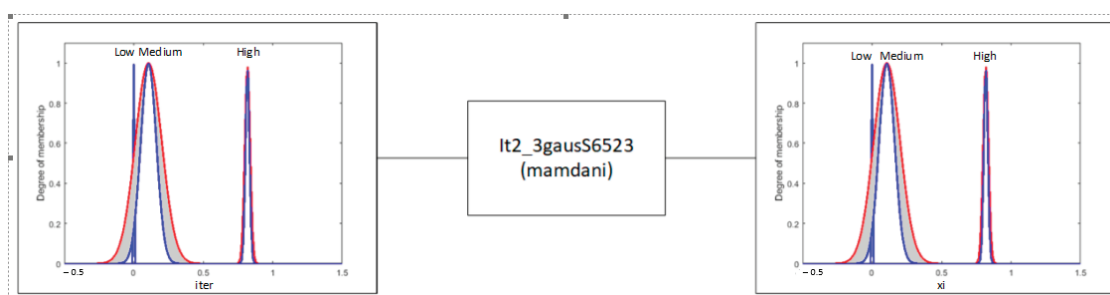


Figure 13. Architecture for IT2FLS FIS-it2_3gausS6523.

Figure 14 shows the Gaussian Membership Function with the uncertainty in the standard deviation that is used in the design of architecture of the interval type-2 fuzzy logic system. Within the framework of the three Gaussian membership functions, namely Low, Medium, and High, the upper membership function (UMF) is represented by the red line, while the lower membership function (LMF) is represented by the blue line. The gray-colored region within the membership function is referred to as the footprint of uncertainty (FOU).

Equations (10)–(13) represent the type-2 Gaussian equations.

$$\mu_{\tilde{F}} = \left[\mu_{\tilde{F}}(x), \mu_{\tilde{F}}(x) = \text{igaussstype2}(x, [\sigma_1, \sigma_2, m]) \right] \quad (10)$$

$$\mu_{\tilde{F}} = \exp \left[-\frac{1}{2} \left(\frac{x-m}{\tilde{\sigma}} \right)^2 \right], \tilde{\sigma} \in [\sigma_1, \sigma_2] \quad (11)$$

$$\mu_{\tilde{F}}(x) = \exp \left[-\frac{1}{2} \left(\frac{x-m}{\sigma_2} \right)^2 \right] \quad (12)$$

$$\mu_{\tilde{F}}(x) = \exp \left[-\frac{1}{2} \left(\frac{x-m}{\sigma_1} \right)^2 \right] \quad (13)$$

Table 4 shows the configuration of the two type-1 and type-2 fuzzy systems that are compared in this investigation.

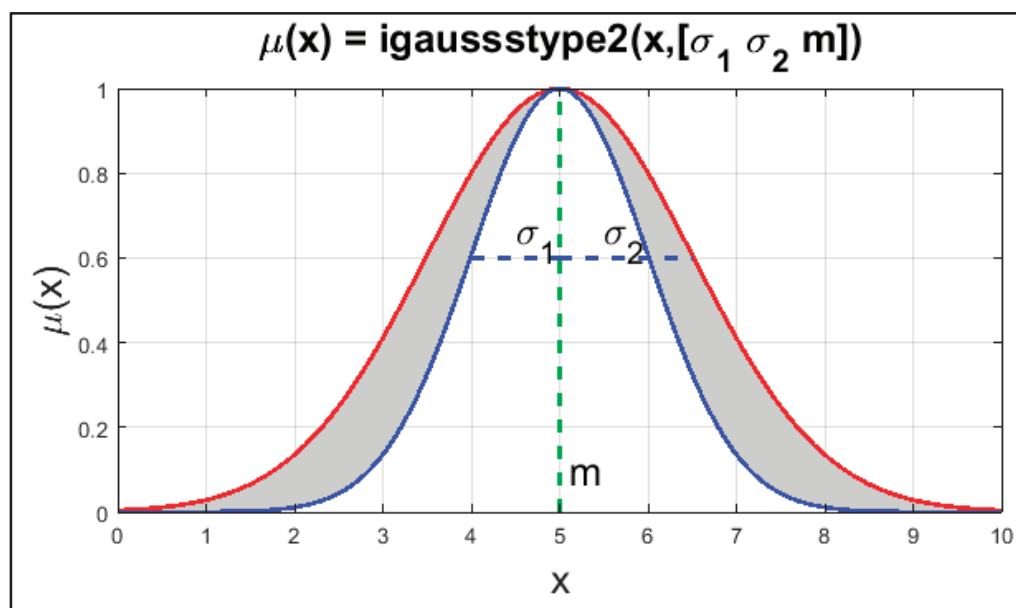


Figure 14. Gaussian Membership Function with uncertainty in the standard deviation.

Table 4. Fuzzy systems configuration T1FLS and IT2FLS.

T1FLS = "fisGau318"	[System] Name = "fisGau318" Type = "mamdani" Version = 2.0 NumInputs = 1 NumOutputs = 1 NumRules = 3 AndMethod = "min" OrMethod = "max" ImpMethod = "min" AggMethod = "max" DefuzzMethod = "centroid"
IT2FLS = "it2_3gausS6523"	[System] Name = "it2_3gausS6523" Type = "mamdani" Version = 2.0 NumInputs = 1 NumOutputs = 1 NumRules = 3 AndMethod = "min" OrMethod = "max" ImpMethod = "min" AggMethod = "max" DefuzzMethod = "centroid"

6. Results

Table 5 shows the mean and standard deviation of T1FLS-fisGau318 for 30, 50, and 100 dimensions, for the 36 mathematical functions in Table 3.

Table 6 shows the mean and SD for IT2FLS-it2_3gausS6523 optimized with the DMOA algorithm for 30, 50, and 100 dimensions.

Comparison Table 7 shows the hypothesis test between T1FLS and IT2FLS fuzzy systems with parameter adaptation and 30 dimensions where the IT2FLS (Not optimized) was better in 5 of 36 mathematical functions.

Table 5. Mean and standard deviation for T1FLS and for 30, 50, and 100 dimensions.

N	T1FLS-DMOAx30		T1FLS-DMOAx50		T1FLS-DMOAx100	
	Mean	SD	Mean	SD	Mean	SD
1	1.04×10^{-12}	2.82×10^{-12}	1.30×10^{-12}	3.49×10^{-12}	3.86×10^{-12}	1.66×10^{-11}
2	2.91×10^{-15}	6.07×10^{-15}	7.96×10^{-16}	2.01×10^{-15}	7.39×10^{-15}	2.63×10^{-14}
3	5.62×10^{-17}	2.12×10^{-16}	1.90×10^{-16}	5.43×10^{-16}	3.01×10^{-14}	8.31×10^{-14}
4	1.64×10^{-12}	4.37×10^{-12}	8.02×10^{-13}	1.60×10^{-12}	8.71×10^{-13}	3.38×10^{-12}
5	2.14×10^{-9}	1.79×10^{-9}	2.48×10^{-9}	1.35×10^{-8}	1.93×10^{-8}	1.30×10^{-8}
6	7.54×10^{-13}	1.88×10^{-12}	1.47×10^{-12}	3.39×10^{-12}	7.90×10^{-13}	2.03×10^{-12}
7	2.63×10^{-13}	4.95×10^{-13}	1.46×10^{-12}	4.47×10^{-12}	2.13×10^{-12}	6.74×10^{-12}
8	2.39×10^{-13}	6.15×10^{-13}	1.45×10^{-12}	4.24×10^{-12}	8.91×10^{-13}	2.71×10^{-12}
9	8.35×10^{-13}	2.23×10^{-12}	1.01×10^{-12}	2.80×10^{-12}	5.46×10^{-13}	1.17×10^{-12}
10	1.14×10^{-12}	3.03×10^{-12}	9.81×10^{-13}	2.42×10^{-12}	9.03×10^{-13}	1.97×10^{-12}
11	5.77×10^{-13}	1.16×10^{-12}	4.46×10^{-13}	7.81×10^{-13}	1.22×10^{-12}	3.91×10^{-12}
12	7.45×10^{-13}	1.58×10^{-12}	6.76×10^{-15}	1.34×10^{-14}	2.71×10^{-20}	7.10×10^{-20}
13	1.15×10^{-12}	3.55×10^{-12}	1.48×10^{-12}	3.89×10^{-12}	2.11×10^{-12}	5.48×10^{-12}
14	4.96×10^{-13}	1.20×10^{-12}	1.99×10^{-12}	7.02×10^{-12}	4.24×10^{-13}	8.07×10^{-13}
15	5.46×10^{-13}	1.12×10^{-12}	4.11×10^{-13}	9.33×10^{-13}	1.28×10^{-13}	2.52×10^{-13}
16	6.03×10^{-13}	1.28×10^{-12}	5.79×10^{-13}	1.11×10^{-12}	5.58×10^{-13}	1.19×10^{-12}
17	1.62×10^{-15}	3.07×10^{-15}	5.12×10^{-16}	8.36×10^{-16}	4.69×10^{-16}	1.66×10^{-15}
18	4.31×10^{-17}	9.02×10^{-17}	1.46×10^{-16}	3.83×10^{-16}	8.38×10^{-17}	2.94×10^{-16}
19	6.38×10^{-13}	1.80×10^{-12}	4.32×10^{-12}	1.65×10^{-11}	9.12×10^{-13}	1.68×10^{-12}
20	2.66×10^{-12}	6.49×10^{-12}	7.51×10^{-13}	1.86×10^{-12}	7.58×10^{-13}	1.62×10^{-12}
21	1.46×10^{-12}	7.82×10^{-12}	4.73×10^{-13}	2.32×10^{-12}	2.98×10^{-14}	6.60×10^{-14}
22	1.47×10^{-12}	2.63×10^{-12}	8.50×10^{-13}	3.52×10^{-12}	1.09×10^{-12}	3.49×10^{-12}
23	1.30×10^{-12}	3.46×10^{-12}	5.72×10^{-13}	1.03×10^{-12}	3.76×10^{-13}	8.72×10^{-13}
24	6.63×10^{-13}	1.39×10^{-12}	3.60×10^{-12}	1.83×10^{-11}	2.26×10^{-12}	7.21×10^{-12}
25	1.65×10^{-15}	5.19×10^{-15}	1.43×10^{-15}	4.36×10^{-15}	6.01×10^{-16}	1.50×10^{-15}
26	2.40×10^{-8}	1.36×10^{-8}	1.79×10^{-8}	1.05×10^{-8}	1.87×10^{-8}	9.88×10^{-9}
27	1.57×10^{-13}	2.55×10^{-13}	3.87×10^{-13}	8.95×10^{-13}	2.11×10^{-12}	7.50×10^{-12}
28	3.95×10^{-13}	9.97×10^{-13}	4.16×10^{-13}	8.78×10^{-13}	4.11×10^{-12}	9.48×10^{-12}
29	2.18×10^{-12}	6.80×10^{-12}	2.97×10^{-13}	6.44×10^{-13}	6.45×10^{-13}	1.02×10^{-12}
30	1.56×10^{-12}	3.32×10^{-12}	4.08×10^{-13}	7.22×10^{-13}	2.99×10^{-13}	5.03×10^{-13}
31	1.61×10^{-12}	3.89×10^{-12}	2.99×10^{-12}	1.13×10^{-11}	9.34×10^{-13}	2.65×10^{-12}
32	1.37×10^{-12}	3.30×10^{-12}	1.67×10^{-13}	4.69×10^{-13}	1.71×10^{-13}	5.44×10^{-13}
33	2.32×10^{-12}	6.87×10^{-12}	1.76×10^{-12}	4.06×10^{-12}	1.33×10^{-12}	3.02×10^{-12}
34	1.09×10^{-20}	3.96×10^{-20}	3.74×10^{-20}	1.16×10^{-19}	5.56×10^{-20}	2.29×10^{-19}
35	1.79×10^{-12}	5.13×10^{-12}	3.83×10^{-13}	7.54×10^{-13}	4.14×10^{-13}	8.86×10^{-13}
36	1.25×10^{-12}	3.25×10^{-12}	6.11×10^{-13}	1.02×10^{-12}	1.55×10^{-12}	4.83×10^{-12}

Table 6. Mean and standard deviation for IT2FLS and for 30, 50, and 100 dimensions.

N	IT2FLS-DMOAx30		IT2FLS-DMOAx50		IT2FLS-DMOAx100	
	Mean	SD	Mean	SD	Mean	SD
1	5.39×10^{-20}	1.18×10^{-19}	3.76×10^{-20}	9.84×10^{-20}	3.97×10^{-20}	6.97×10^{-20}
2	9.69×10^{-22}	2.04×10^{-21}	7.56×10^{-22}	9.99×10^{-22}	7.99×10^{-22}	1.81×10^{-21}
3	1.02×10^{-22}	2.36×10^{-22}	1.81×10^{-22}	4.36×10^{-22}	6.00×10^{-21}	1.10×10^{-20}
4	3.27×10^{-20}	6.48×10^{-20}	4.19×10^{-20}	8.97×10^{-20}	6.70×10^{-20}	1.42×10^{-19}
5	1.79×10^{-18}	1.30×10^{-18}	2.80×10^{-18}	5.28×10^{-18}	6.42×10^{-18}	4.23×10^{-18}
6	1.75×10^{-20}	4.76×10^{-20}	2.80×10^{-20}	3.96×10^{-20}	5.60×10^{-20}	7.87×10^{-20}
7	2.57×10^{-20}	3.54×10^{-20}	5.97×10^{-20}	1.20×10^{-19}	6.67×10^{-20}	2.76×10^{-19}
8	2.08×10^{-20}	4.94×10^{-20}	7.59×10^{-20}	2.55×10^{-19}	5.84×10^{-20}	1.36×10^{-19}
9	2.51×10^{-20}	6.85×10^{-20}	4.39×10^{-20}	7.03×10^{-20}	3.56×10^{-20}	6.70×10^{-20}
10	3.24×10^{-20}	6.93×10^{-20}	4.86×10^{-20}	8.13×10^{-20}	3.56×10^{-20}	6.07×10^{-20}
11	8.33×10^{-20}	1.64×10^{-19}	4.14×10^{-20}	6.42×10^{-20}	6.20×10^{-20}	8.90×10^{-20}
12	2.89×10^{-20}	4.59×10^{-20}	1.64×10^{-20}	5.89×10^{-20}	6.23×10^{-24}	4.84×10^{-24}
13	1.25×10^{-20}	1.78×10^{-20}	5.64×10^{-20}	9.65×10^{-20}	1.70×10^{-20}	2.93×10^{-20}
14	4.10×10^{-20}	7.01×10^{-20}	2.63×10^{-20}	4.20×10^{-20}	3.15×10^{-20}	6.48×10^{-20}
15	3.07×10^{-20}	4.41×10^{-20}	2.63×10^{-20}	4.71×10^{-20}	4.77×10^{-20}	8.87×10^{-20}
16	2.01×10^{-20}	5.14×10^{-20}	2.11×10^{-20}	2.65×10^{-20}	2.86×10^{-20}	3.96×10^{-20}
17	7.19×10^{-22}	1.03×10^{-21}	1.06×10^{-21}	1.78×10^{-21}	8.02×10^{-22}	1.53×10^{-21}

Table 6. Cont.

N	IT2FLS-DMOAx30		IT2FLS-DMOAx50		IT2FLS-DMOAx100	
	Mean	SD	Mean	SD	Mean	SD
18	2.19×10^{-22}	5.72×10^{-22}	1.78×10^{-22}	7.21×10^{-22}	5.96×10^{-23}	1.08×10^{-22}
19	3.19×10^{-20}	6.17×10^{-20}	2.90×10^{-20}	5.45×10^{-20}	3.73×10^{-20}	4.39×10^{-20}
20	3.29×10^{-20}	5.28×10^{-20}	3.47×10^{-20}	5.03×10^{-20}	2.33×10^{-20}	3.70×10^{-20}
21	1.37×10^{-20}	4.70×10^{-20}	6.02×10^{-21}	1.12×10^{-20}	1.78×10^{-20}	4.97×10^{-20}
22	4.76×10^{-20}	1.54×10^{-19}	6.65×10^{-20}	1.92×10^{-19}	1.49×10^{-20}	2.85×10^{-20}
23	1.82×10^{-20}	3.14×10^{-20}	1.05×10^{-19}	1.74×10^{-19}	1.84×10^{-20}	3.05×10^{-20}
24	3.25×10^{-20}	6.89×10^{-20}	3.34×10^{-20}	6.81×10^{-20}	2.24×10^{-20}	4.26×10^{-20}
25	9.69×10^{-22}	2.29×10^{-21}	9.72×10^{-22}	2.36×10^{-21}	3.30×10^{-22}	6.23×10^{-22}
26	6.05×10^{-18}	4.05×10^{-18}	7.26×10^{-18}	5.36×10^{-18}	7.70×10^{-18}	5.55×10^{-18}
27	7.03×10^{-20}	1.14×10^{-19}	4.36×10^{-20}	8.94×10^{-20}	1.53×10^{-20}	2.38×10^{-20}
28	5.67×10^{-20}	1.09×10^{-19}	4.68×10^{-20}	1.01×10^{-19}	4.93×10^{-20}	8.81×10^{-20}
29	4.76×10^{-20}	9.85×10^{-20}	3.95×10^{-20}	6.92×10^{-20}	3.25×10^{-20}	5.28×10^{-20}
30	3.38×10^{-20}	6.08×10^{-20}	3.06×10^{-20}	6.69×10^{-20}	2.19×10^{-20}	4.23×10^{-20}
31	6.30×10^{-20}	7.13×10^{-20}	5.60×10^{-20}	1.23×10^{-19}	2.12×10^{-20}	3.97×10^{-20}
32	2.71×10^{-20}	6.28×10^{-20}	4.81×10^{-20}	9.99×10^{-20}	2.56×10^{-20}	4.14×10^{-20}
33	4.24×10^{-20}	8.03×10^{-20}	2.74×10^{-20}	4.33×10^{-20}	2.55×10^{-20}	3.25×10^{-20}
34	4.68×10^{-24}	2.31×10^{-24}	6.08×10^{-24}	3.46×10^{-24}	4.96×10^{-24}	2.05×10^{-24}
35	2.81×10^{-20}	4.92×10^{-20}	2.64×10^{-20}	3.01×10^{-20}	1.23×10^{-19}	3.19×10^{-19}
36	9.93×10^{-20}	2.71×10^{-19}	6.27×10^{-20}	1.32×10^{-19}	2.53×10^{-20}	4.93×10^{-20}

Table 7. Hypothesis Test T1FLS vs. IT2FLS—30 dimensions.

No	T1FLS		IT2FLS		Hypothesis Test	
	fisGau318 30		it2_3gausS01 30			
	Mean	SD	Mean	SD	Z	E
1	1.04×10^{-12}	2.82×10^{-12}	1.44×10^{-12}	4.82×10^{-12}	-0.94	N
2	2.91×10^{-15}	6.07×10^{-15}	2.93×10^{-15}	7.79×10^{-15}	-2.08	Y
3	5.62×10^{-17}	2.12×10^{-16}	5.54×10^{-16}	2.08×10^{-15}	1.96	N
4	1.64×10^{-12}	4.37×10^{-12}	2.83×10^{-12}	1.09×10^{-11}	-0.99	N
5	2.14×10^{-9}	1.79×10^{-9}	1.14×10^{-9}	1.08×10^{-9}	-2.8	Y
6	7.54×10^{-13}	1.88×10^{-12}	2.47×10^{-13}	5.68×10^{-13}	-0.04	N
7	2.63×10^{-13}	4.95×10^{-13}	1.46×10^{-12}	5.09×10^{-12}	1.62	N
8	2.39×10^{-13}	6.15×10^{-13}	5.53×10^{-13}	2.02×10^{-12}	0.77	N
9	8.35×10^{-13}	2.23×10^{-12}	4.85×10^{-13}	1.16×10^{-12}	-0.35	N
10	1.14×10^{-12}	3.03×10^{-12}	3.49×10^{-13}	6.59×10^{-13}	0.09	N
11	5.77×10^{-13}	1.16×10^{-12}	1.20×10^{-12}	3.45×10^{-12}	0.29	N
12	7.45×10^{-13}	1.58×10^{-12}	9.57×10^{-13}	3.18×10^{-12}	-0.31	N
13	1.15×10^{-12}	3.55×10^{-12}	6.23×10^{-13}	2.15×10^{-12}	-0.84	N
14	4.96×10^{-13}	1.20×10^{-12}	2.39×10^{-13}	4.22×10^{-13}	1.54	N
15	5.46×10^{-13}	1.12×10^{-12}	2.30×10^{-13}	3.98×10^{-13}	1.15	N
16	6.03×10^{-13}	1.28×10^{-12}	8.31×10^{-13}	2.34×10^{-12}	0.92	N
17	1.62×10^{-15}	3.07×10^{-15}	3.09×10^{-15}	1.38×10^{-14}	0.63	N
18	4.31×10^{-17}	9.02×10^{-17}	9.31×10^{-16}	3.65×10^{-15}	1.47	N
19	6.38×10^{-13}	1.80×10^{-12}	3.87×10^{-13}	7.05×10^{-13}	-1.15	N
20	2.66×10^{-12}	6.49×10^{-12}	3.41×10^{-13}	6.03×10^{-13}	-1.1	N
21	1.46×10^{-12}	7.82×10^{-12}	1.97×10^{-14}	3.75×10^{-14}	-0.96	N
22	1.47×10^{-12}	2.63×10^{-12}	2.22×10^{-12}	9.69×10^{-12}	-2.15	Y
23	1.30×10^{-12}	3.46×10^{-12}	8.46×10^{-13}	2.37×10^{-12}	0.78	N
24	6.63×10^{-13}	1.39×10^{-12}	9.68×10^{-13}	3.02×10^{-12}	0.69	N
25	1.65×10^{-15}	5.19×10^{-15}	1.79×10^{-15}	7.18×10^{-15}	-0.68	N
26	2.40×10^{-8}	1.36×10^{-8}	1.44×10^{-8}	7.47×10^{-9}	-3.81	Y
27	1.57×10^{-13}	2.55×10^{-13}	3.18×10^{-13}	7.86×10^{-13}	-0.24	N
28	3.95×10^{-13}	9.97×10^{-13}	1.13×10^{-13}	1.47×10^{-13}	-0.66	N
29	2.18×10^{-12}	6.80×10^{-12}	2.05×10^{-12}	5.11×10^{-12}	-1.61	N
30	1.56×10^{-12}	3.32×10^{-12}	1.37×10^{-12}	3.86×10^{-12}	-1.97	Y
31	1.61×10^{-12}	3.89×10^{-12}	2.26×10^{-13}	4.34×10^{-13}	-0.07	N
32	1.37×10^{-12}	3.30×10^{-12}	4.40×10^{-13}	1.22×10^{-12}	-1.25	N
33	2.32×10^{-12}	6.87×10^{-12}	1.45×10^{-12}	6.94×10^{-12}	-0.97	N
34	1.09×10^{-20}	3.96×10^{-20}	2.25×10^{-20}	6.20×10^{-20}	0.98	N
35	1.79×10^{-12}	5.13×10^{-12}	3.85×10^{-13}	5.53×10^{-13}	-1.1	N
36	1.25×10^{-12}	3.25×10^{-12}	4.75×10^{-13}	1.30×10^{-12}	-1.31	N

Figures 15 and 16 show the behaviors of the mean and standard deviation of the two types of fuzzy systems T1FLS vs. IT2FLS for 30 dimensions.

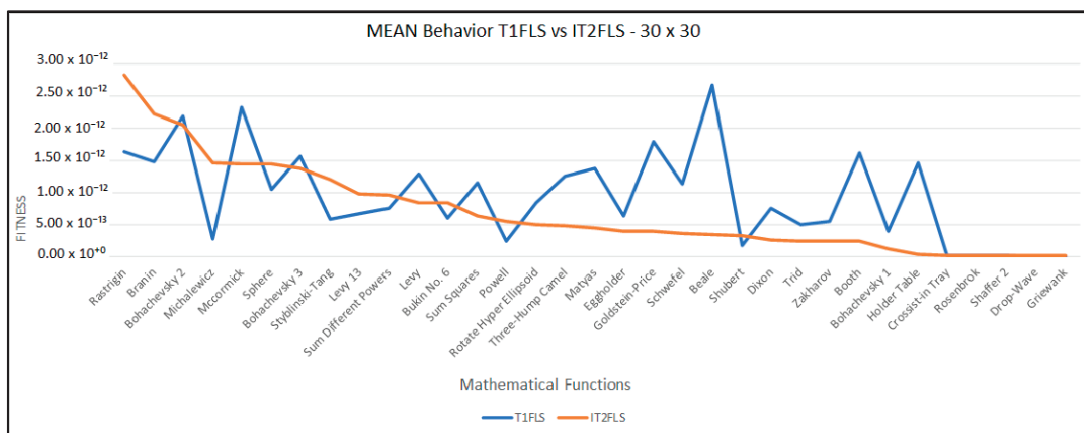


Figure 15. Behavior of the MEAN for T1FLS vs. IT2FLS (Not optimized) 30 dimensions.

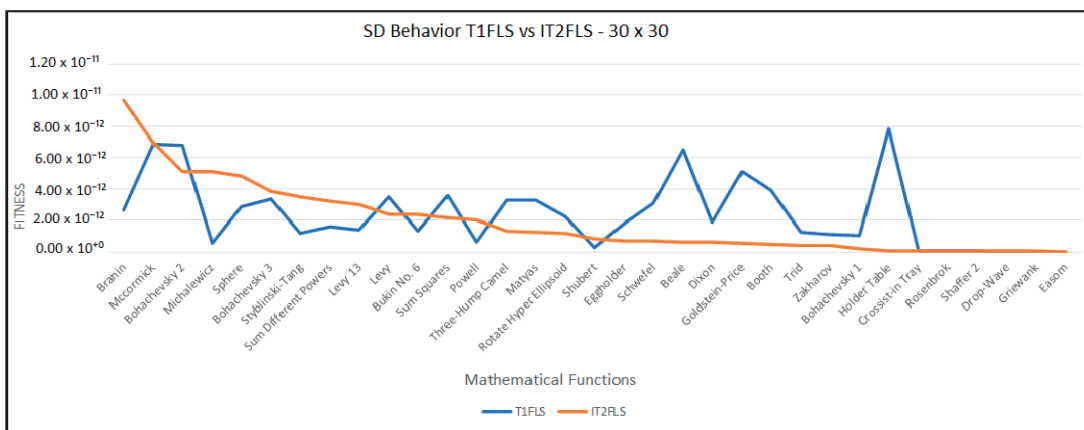


Figure 16. Behavior of the SD for T1FLS vs. IT2FLS (Not optimized) 30 dimensions.

Table 8 shows the convergence of the T1FDMOA method for 30 dimensions of the mathematical Rosenbrock, Griewank, Rastrigin, Ackley, and Dixon functions.

Table 8. Convergence of T1FDMOA (T1FLS).

T1FDMOA					
N	Rosenbrock	Griewank	Rastrigin	Ackley	Dixon
1	2.36×10^{-14}	1.16×10^{-15}	2.02×10^{-11}	6.60×10^{-9}	8.45×10^{-12}
2	1.70×10^{-14}	1.84×10^{-16}	1.30×10^{-11}	6.38×10^{-9}	6.47×10^{-12}
3	1.43×10^{-14}	1.32×10^{-16}	6.56×10^{-12}	5.63×10^{-9}	1.57×10^{-12}
4	1.39×10^{-14}	7.92×10^{-17}	2.57×10^{-12}	5.20×10^{-9}	8.69×10^{-13}
5	7.44×10^{-15}	6.08×10^{-17}	1.40×10^{-12}	4.58×10^{-9}	7.23×10^{-13}
6	3.49×10^{-15}	1.79×10^{-17}	1.11×10^{-12}	3.00×10^{-9}	6.88×10^{-13}
7	2.25×10^{-15}	1.35×10^{-17}	8.16×10^{-13}	2.82×10^{-9}	6.56×10^{-13}
8	2.07×10^{-15}	1.13×10^{-17}	7.81×10^{-13}	2.52×10^{-9}	5.66×10^{-13}
9	1.23×10^{-15}	8.50×10^{-18}	6.32×10^{-13}	2.50×10^{-9}	5.20×10^{-13}
10	4.85×10^{-16}	6.81×10^{-18}	6.01×10^{-13}	2.05×10^{-9}	5.07×10^{-13}
11	3.92×10^{-16}	4.42×10^{-18}	3.02×10^{-13}	1.99×10^{-9}	3.95×10^{-13}
12	3.77×10^{-16}	3.11×10^{-18}	2.84×10^{-13}	1.96×10^{-9}	3.30×10^{-13}
13	2.16×10^{-16}	1.56×10^{-18}	2.63×10^{-13}	1.94×10^{-9}	1.80×10^{-13}

Table 8. Cont.

T1DMOA					
N	Rosenbrock	Griewank	Rastrigin	Ackley	Dixon
14	1.38×10^{-16}	1.19×10^{-18}	2.48×10^{-13}	1.78×10^{-9}	1.63×10^{-13}
15	1.07×10^{-16}	8.77×10^{-19}	9.11×10^{-14}	1.70×10^{-9}	9.61×10^{-14}
16	8.13×10^{-17}	7.96×10^{-19}	7.89×10^{-14}	1.52×10^{-9}	9.49×10^{-14}
17	8.07×10^{-17}	5.67×10^{-19}	6.61×10^{-14}	1.41×10^{-9}	9.39×10^{-14}
18	7.83×10^{-17}	3.79×10^{-19}	4.30×10^{-14}	1.32×10^{-9}	8.95×10^{-14}
19	6.97×10^{-17}	3.46×10^{-19}	2.72×10^{-14}	1.21×10^{-9}	5.51×10^{-14}
20	6.69×10^{-17}	2.09×10^{-19}	2.01×10^{-14}	1.14×10^{-9}	5.21×10^{-14}
21	3.75×10^{-17}	1.90×10^{-19}	2.01×10^{-14}	1.10×10^{-9}	1.50×10^{-14}
22	1.67×10^{-17}	1.63×10^{-19}	8.26×10^{-15}	1.01×10^{-9}	1.12×10^{-14}
23	1.63×10^{-17}	1.62×10^{-19}	3.35×10^{-15}	9.97×10^{-10}	9.05×10^{-15}
24	5.22×10^{-18}	1.09×10^{-19}	2.88×10^{-15}	9.87×10^{-10}	8.08×10^{-15}
25	1.13×10^{-19}	5.59×10^{-20}	2.81×10^{-15}	9.40×10^{-10}	7.59×10^{-15}
26	9.94×10^{-20}	4.94×10^{-20}	7.21×10^{-16}	6.83×10^{-10}	1.19×10^{-16}
27	4.94×10^{-20}	1.37×10^{-20}	5.87×10^{-16}	5.36×10^{-10}	8.47×10^{-17}
28	2.38×10^{-20}	5.92×10^{-21}	3.10×10^{-16}	5.13×10^{-10}	3.28×10^{-17}
29	5.52×10^{-22}	2.01×10^{-21}	6.27×10^{-18}	2.08×10^{-10}	2.82×10^{-18}
30	2.18×10^{-22}	1.90×10^{-21}	9.69×10^{-21}	2.35×10^{-16}	5.76×10^{-20}

Table 9 shows the convergence of the IT2FDMOA method for 30 dimensions of the mathematical Rosenbrock, Griewank, Rastrigin, Ackley, and Dixon function.

Table 9. Convergence of IT2FDMOA (IT2FLS).

IT2DMOA					
N	Rosenbrock	Griewank	Rastrigin	Ackley	Dixon
1	1.10×10^{-20}	1.15×10^{-21}	3.30×10^{-19}	7.33×10^{-18}	2.56×10^{-19}
2	2.53×10^{-21}	6.51×10^{-22}	1.50×10^{-19}	3.44×10^{-18}	7.39×10^{-20}
3	2.14×10^{-21}	2.49×10^{-22}	8.57×10^{-20}	3.18×10^{-18}	3.87×10^{-20}
4	1.89×10^{-21}	1.92×10^{-22}	7.09×10^{-20}	2.91×10^{-18}	3.20×10^{-20}
5	1.73×10^{-21}	1.79×10^{-22}	5.33×10^{-20}	2.70×10^{-18}	2.72×10^{-20}
6	1.60×10^{-21}	1.40×10^{-22}	4.24×10^{-20}	2.56×10^{-18}	1.75×10^{-20}
7	1.57×10^{-21}	7.15×10^{-23}	4.15×10^{-20}	2.41×10^{-18}	1.53×10^{-20}
8	1.33×10^{-21}	5.31×10^{-23}	3.40×10^{-20}	2.07×10^{-18}	1.00×10^{-20}
9	1.29×10^{-21}	4.75×10^{-23}	2.41×10^{-20}	2.03×10^{-18}	8.16×10^{-21}
10	9.88×10^{-22}	3.73×10^{-23}	2.26×10^{-20}	1.86×10^{-18}	7.74×10^{-21}
11	7.09×10^{-22}	3.61×10^{-23}	2.03×10^{-20}	1.81×10^{-18}	6.86×10^{-21}
12	6.90×10^{-22}	3.13×10^{-23}	1.34×10^{-20}	1.62×10^{-18}	5.37×10^{-21}
13	4.46×10^{-22}	2.70×10^{-23}	1.16×10^{-20}	1.55×10^{-18}	4.83×10^{-21}
14	2.83×10^{-22}	2.62×10^{-23}	1.08×10^{-20}	1.48×10^{-18}	4.62×10^{-21}
15	2.25×10^{-22}	2.54×10^{-23}	1.04×10^{-20}	1.45×10^{-18}	4.03×10^{-21}
16	1.82×10^{-22}	2.31×10^{-23}	9.73×10^{-21}	1.43×10^{-18}	3.72×10^{-21}
17	1.47×10^{-22}	2.04×10^{-23}	8.79×10^{-21}	1.41×10^{-18}	2.84×10^{-21}
18	9.52×10^{-23}	1.97×10^{-23}	6.98×10^{-21}	1.35×10^{-18}	2.63×10^{-21}
19	8.76×10^{-23}	1.95×10^{-23}	6.55×10^{-21}	1.33×10^{-18}	1.62×10^{-21}
20	5.84×10^{-23}	1.23×10^{-23}	5.24×10^{-21}	1.26×10^{-18}	1.16×10^{-21}
21	3.47×10^{-23}	7.12×10^{-24}	4.86×10^{-21}	1.26×10^{-18}	1.14×10^{-21}
22	3.12×10^{-23}	5.57×10^{-24}	4.16×10^{-21}	1.13×10^{-18}	5.58×10^{-22}
23	3.05×10^{-23}	5.30×10^{-24}	3.13×10^{-21}	1.05×10^{-18}	2.14×10^{-22}
24	9.22×10^{-24}	4.56×10^{-24}	3.07×10^{-21}	1.03×10^{-18}	1.80×10^{-22}
25	5.51×10^{-24}	3.73×10^{-24}	2.86×10^{-21}	9.60×10^{-19}	8.63×10^{-23}
26	4.24×10^{-24}	3.27×10^{-24}	2.07×10^{-21}	8.09×10^{-19}	7.28×10^{-23}

Table 9. Cont.

N	IT2DMOA				
	Rosenbrock	Griewank	Rastrigin	Ackley	Dixon
27	2.44×10^{-24}	2.50×10^{-24}	1.64×10^{-21}	7.61×10^{-19}	3.99×10^{-23}
28	3.98×10^{-25}	2.49×10^{-24}	5.63×10^{-22}	7.07×10^{-19}	2.67×10^{-23}
29	1.39×10^{-25}	1.04×10^{-24}	2.11×10^{-22}	6.77×10^{-19}	1.28×10^{-23}
30	4.22×10^{-26}	1.17×10^{-25}	1.96×10^{-22}	1.88×10^{-19}	6.05×10^{-24}

In Figure 17, we can see how the two methods T1FDMOA and IT2FDMOA converge for the Rosenbrock mathematical function, in Figures 18 and 19 we can see separately the two methods so that we can clearly observe the convergence and the exploration and exploitation phases of each of them, both methods tend to 0 but never reach the value of the same mathematical function.

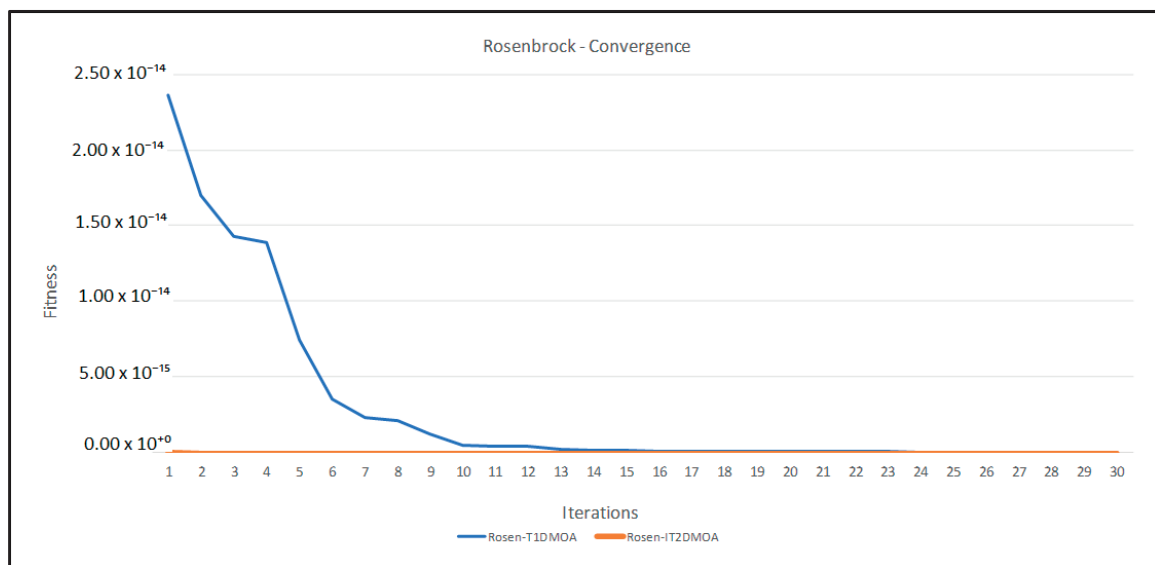


Figure 17. Convergence of the T1FDMOA, and IT2FDMOA methods, for the mathematical Rosenbrock function.

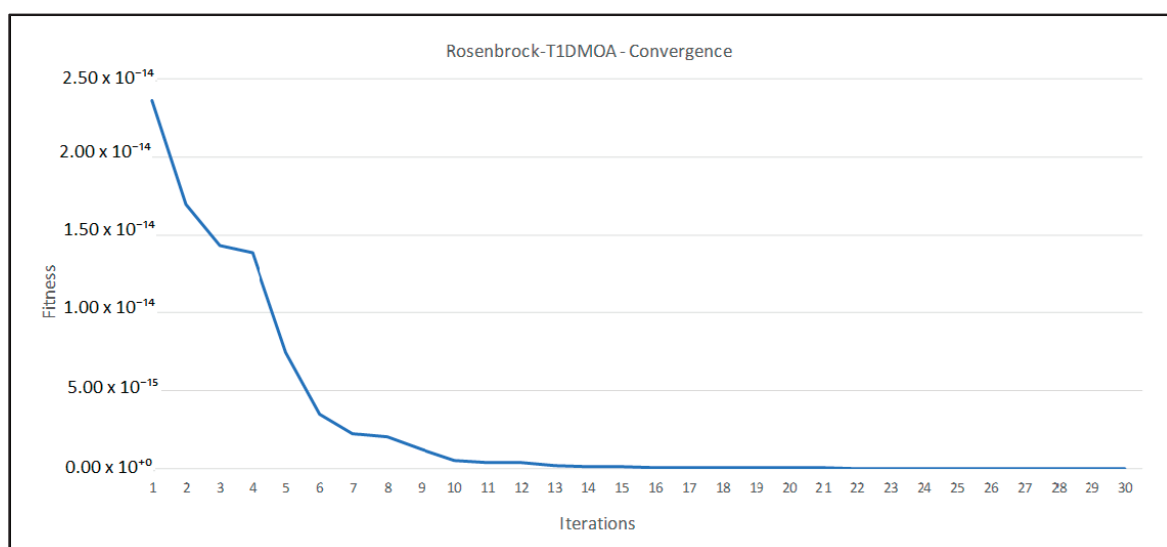


Figure 18. Convergence of the T1FDMOA method, for the mathematical Rosenbrock function.

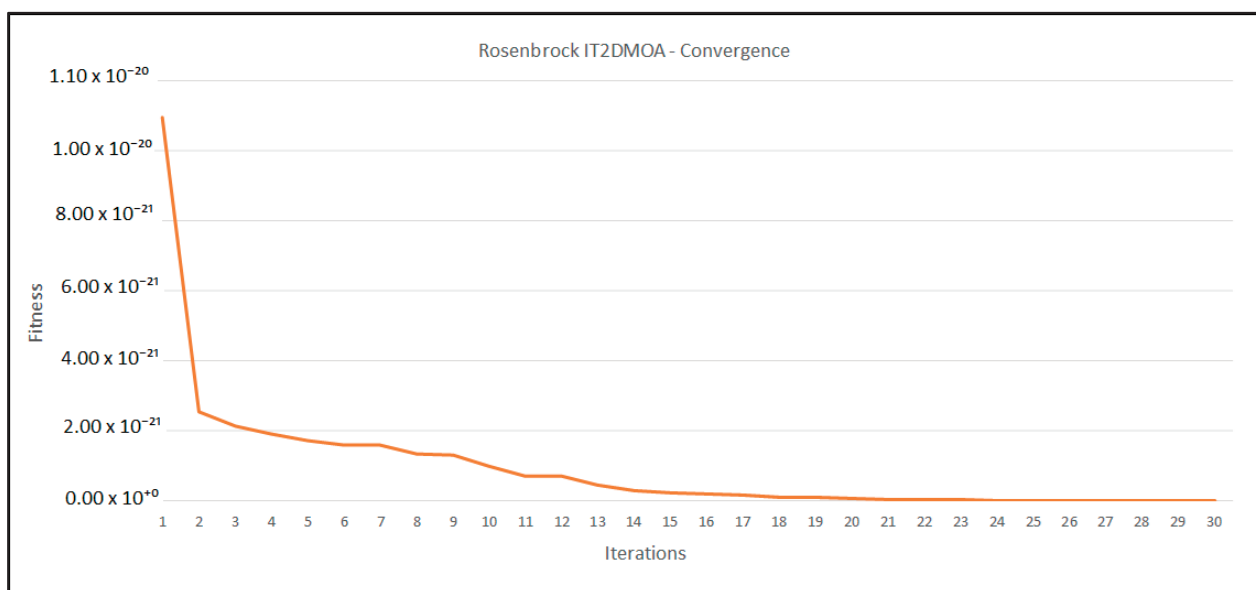


Figure 19. Convergence of the IT2FDMOA method, for the mathematical Rosenbrock function.

Figure 20 illustrates the convergence of the T1FDMOA and IT2FDMOA methods for the Dixon–Price mathematical function. Figures 21 and 22 present a separate visualization of each method, enabling a clear observation of their convergence as well as the exploration and exploitation phases. Although both methods approach but do not reach a value of 0, this behavior is consistent with the characteristics of the Dixon–Price function.

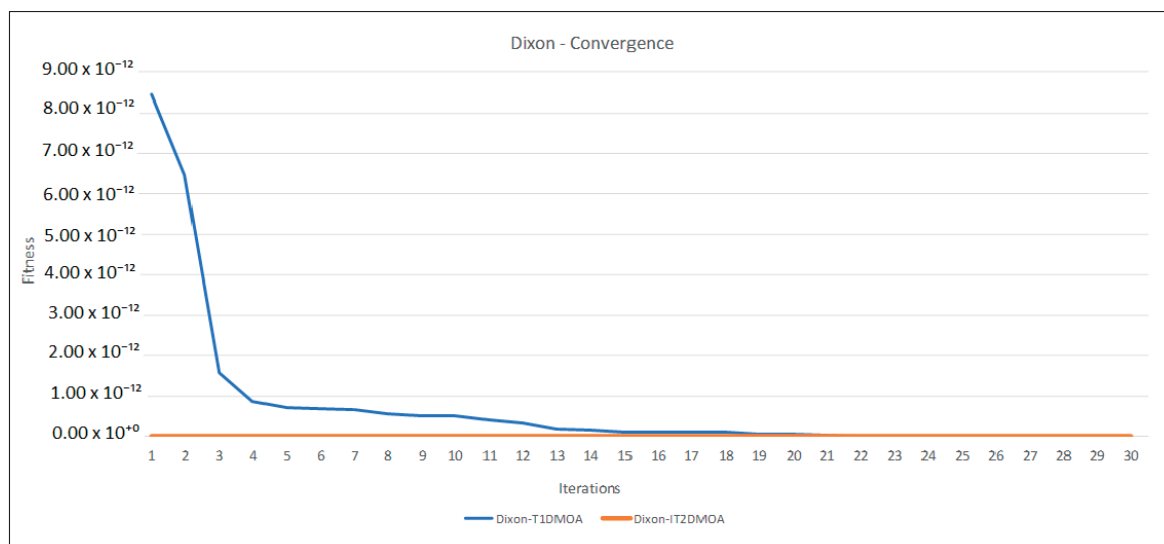


Figure 20. Convergence of the T1FDMOA, and IT2FDMOA methods, for the mathematical Dixon function.

Figure 23 showcases the convergence of the T1FDMOA and IT2FDMOA methods for the Ackley mathematical function. Figures 24 and 25 provide individual depictions of each method, facilitating a distinct analysis of their convergence and the exploration and exploitation phases. As with the previous case, both methods approach but do not reach a value of 0, which aligns with the nature of the Ackley function.

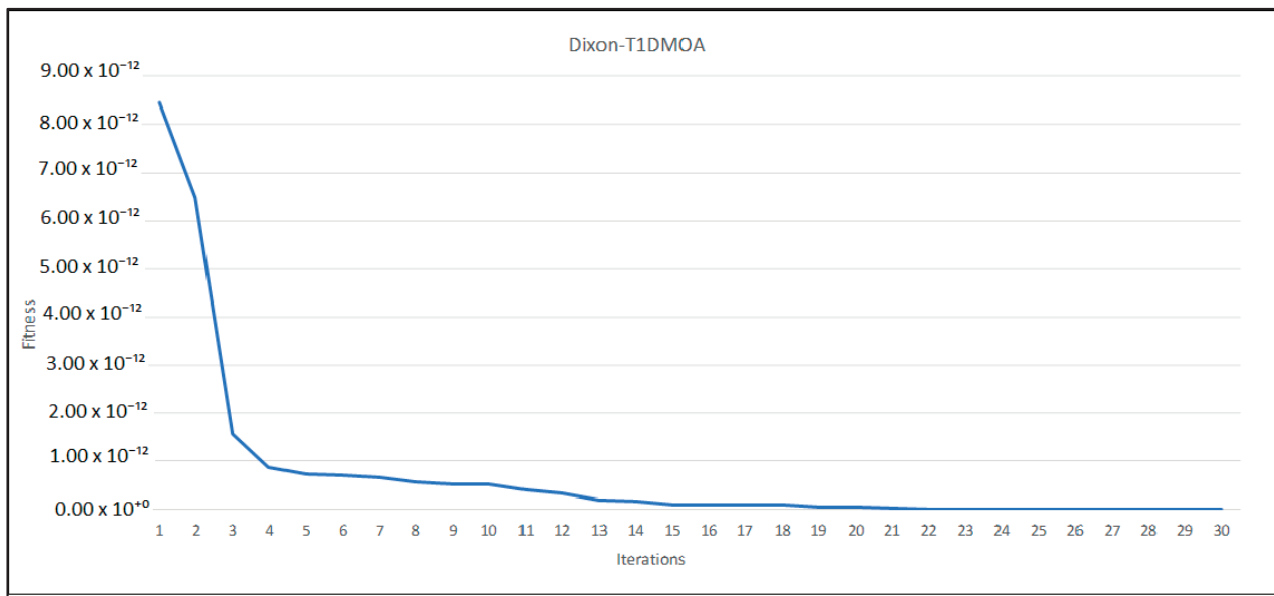


Figure 21. Convergence of the T1FDMOA method, for the mathematical Dixon function.

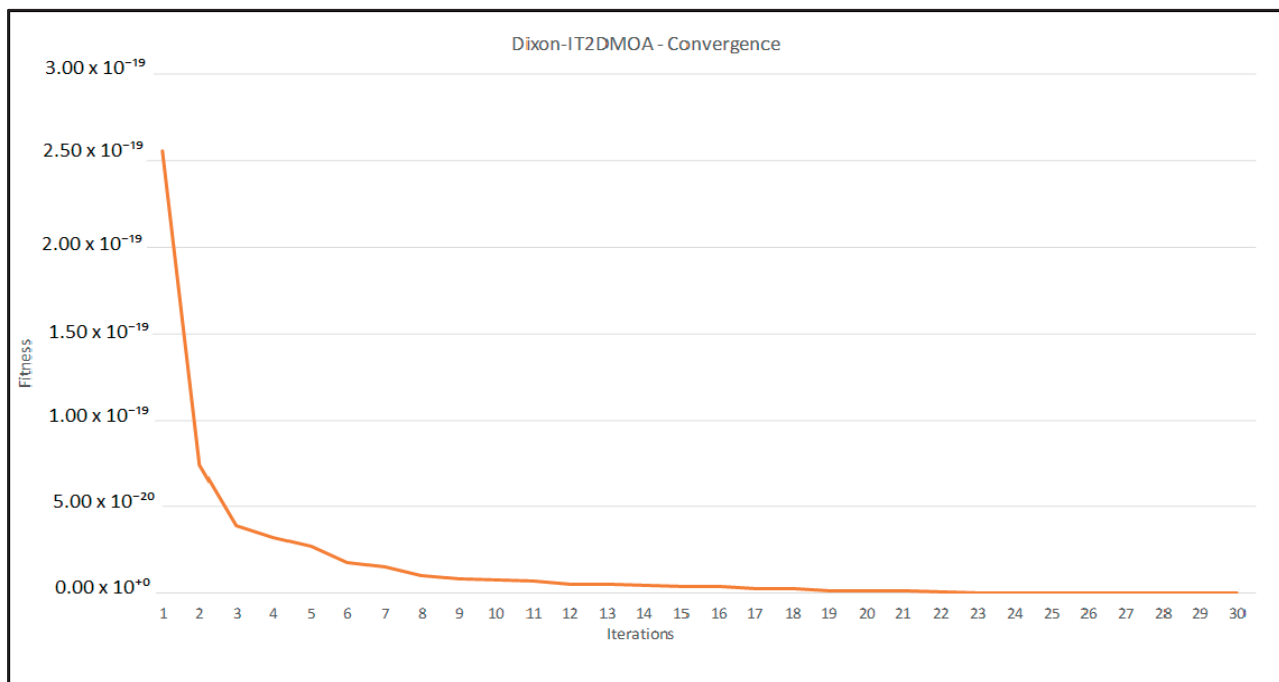


Figure 22. Convergence of the IT2FDMOA method, for the mathematical Dixon function.

6.1. Hypothesis Test

Equation (14) represents the hypothesis test for two independent samples of 30 experiments, the Null Hypothesis Equation (15) and the Alternate Hypothesis Equation (16), with which comparisons were performed between DMOA-T1FLS and DMOA-IT2FLS, where our claim is that the DMOA-IT2FLS method is better than the DMOA-T1FLS method, Figure 26 shows the left-tailed hypothesis test plot.

$$z = \frac{(\bar{x}_1 - \bar{x}_2)}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} \quad (14)$$

$$H_0 : \mu_1 \geq \mu_2 \quad (15)$$

$$H_a : \mu_1 < \mu_2 \text{ claim} \quad (16)$$

Where $\bar{x}_1$ is the mean of sample 1, $\bar{x}_2$ mean of sample 2, σ_1 standard deviation of sample 1, σ_2 standard deviation of sample 2, n_1 number of sample data 1, n_2 number of sample data 2.

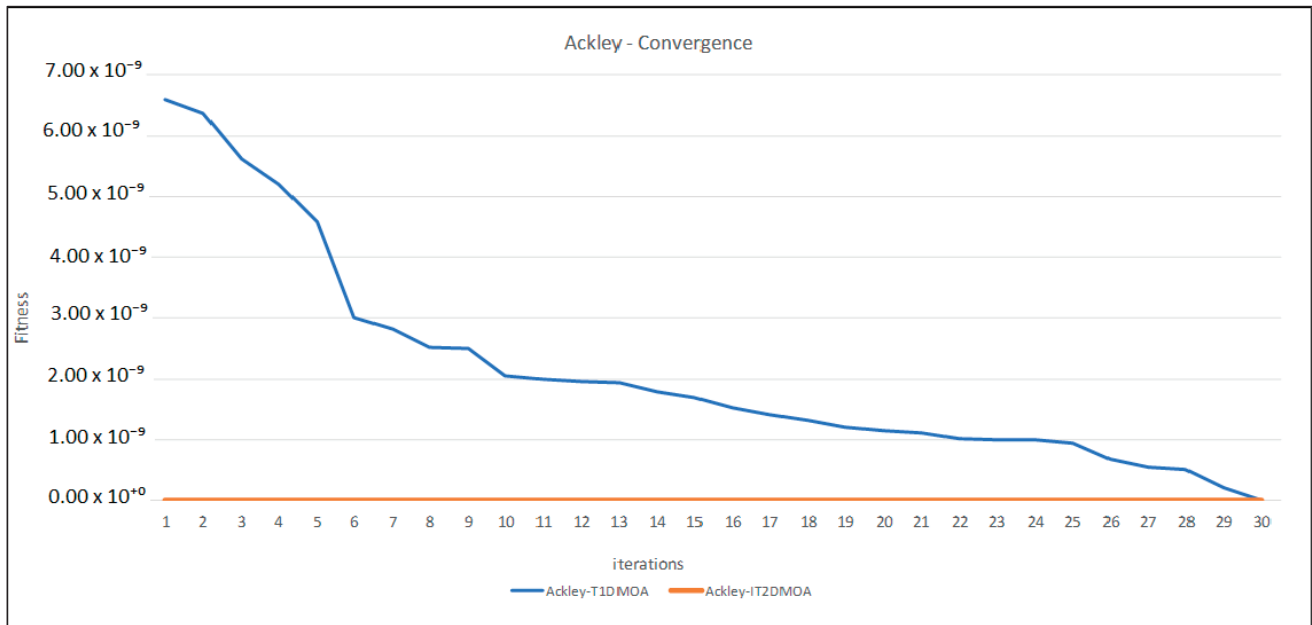


Figure 23. Convergence of the T1FDMOA, and IT2FDMOA methods, for the mathematical Ackley function.

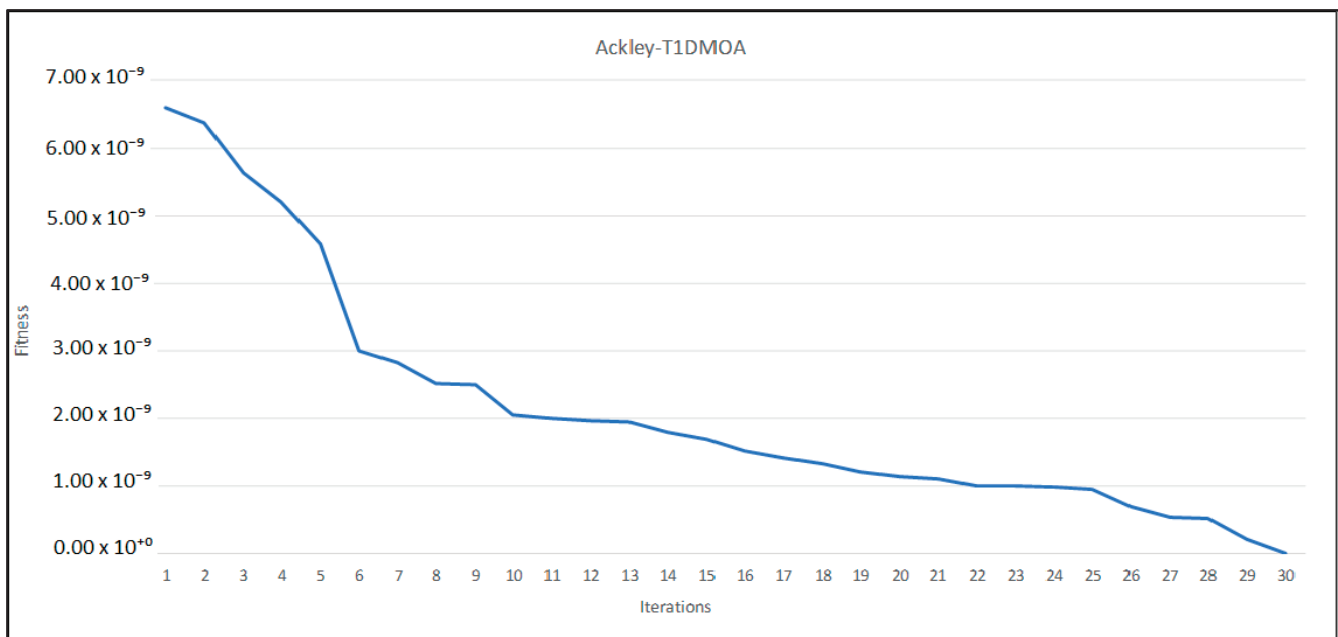


Figure 24. Convergence of the T1FDMOA method, for the mathematical Ackley function.

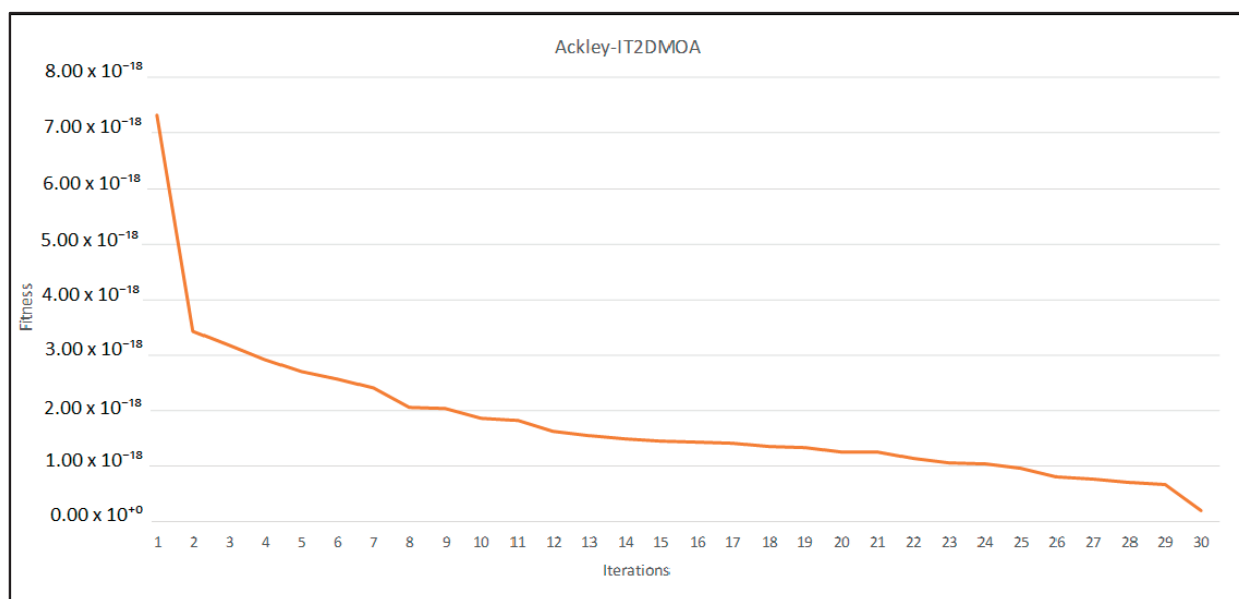


Figure 25. Convergence of the IT2FDMOA method, for the mathematical Ackley function.

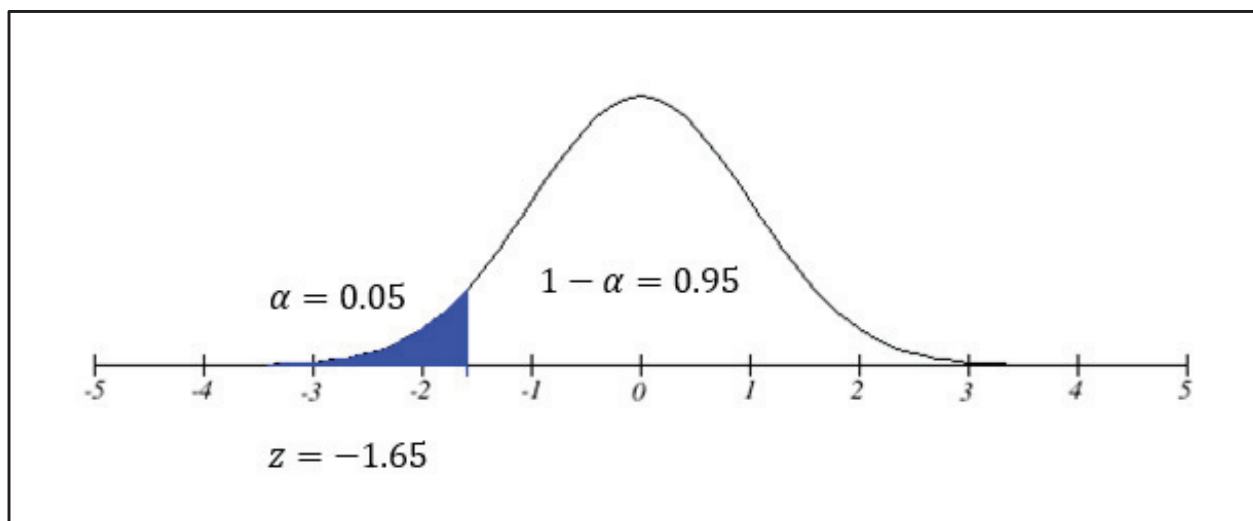


Figure 26. Left-tailed hypothesis test graph.

Significance level $\alpha = 0.05$, confidence level = 95%, confidence level = $1 - \alpha$; $1 - 0.05 = 0.95$ or 95%, since the p -value is less than 0.05, the null hypothesis is rejected.

Comparison Table 10 shows the hypothesis test for 30 dimensions of the fuzzy systems T1FLS with parameter adaptation and optimized IT2FLS was better in 33 out of 36 mathematical functions. In the hypothesis tests in Tables 10–12, our claim is that the results of the experiments performed with optimized IT2FLS are better than the experiments performed with T1FLS with parameter adaptation.

Table 10. Hypothesis test T1FLS vs. IT2FLS—30 dimensions.

No	T1FLS		IT2FLS		Hypothesis Test	
	fisGau318 30		it2_3gausS6523 30			
	Mean	SD	Mean	SD	Z	E
1	1.04×10^{-12}	2.82×10^{-12}	5.39×10^{-20}	1.18×10^{-19}	−2.12	Y
2	2.91×10^{-15}	6.07×10^{-15}	9.69×10^{-22}	2.04×10^{-21}	−2.76	Y

Table 10. Cont.

No	T1FLS		IT2FLS		Hypothesis Test	
	fisGau318 30		it2_3gausS6523 30			
	Mean	SD	Mean	SD	Z	E
3	5.62×10^{-17}	2.12×10^{-16}	1.02×10^{-22}	2.36×10^{-22}	−1.52	N
4	1.64×10^{-12}	4.37×10^{-12}	3.27×10^{-20}	6.48×10^{-20}	−2.16	Y
5	2.14×10^{-9}	1.79×10^{-9}	1.79×10^{-18}	1.30×10^{-18}	−6.88	Y
6	7.54×10^{-13}	1.88×10^{-12}	1.75×10^{-20}	4.76×10^{-20}	−2.31	Y
7	2.63×10^{-13}	4.95×10^{-13}	2.57×10^{-20}	3.54×10^{-20}	−3.06	Y
8	2.39×10^{-13}	6.15×10^{-13}	2.08×10^{-20}	4.94×10^{-20}	−2.24	Y
9	8.35×10^{-13}	2.23×10^{-12}	2.51×10^{-20}	6.85×10^{-20}	−2.16	Y
10	1.14×10^{-12}	3.03×10^{-12}	3.24×10^{-20}	6.93×10^{-20}	−2.16	Y
11	5.77×10^{-13}	1.16×10^{-12}	8.33×10^{-20}	1.64×10^{-19}	−2.86	Y
12	7.45×10^{-13}	1.58×10^{-12}	2.89×10^{-20}	4.59×10^{-20}	−2.72	Y
13	1.15×10^{-12}	3.55×10^{-12}	1.25×10^{-20}	1.78×10^{-20}	−1.86	Y
14	4.96×10^{-13}	1.20×10^{-12}	4.10×10^{-20}	7.01×10^{-20}	−2.37	Y
15	5.46×10^{-13}	1.12×10^{-12}	3.07×10^{-20}	4.41×10^{-20}	−2.8	Y
16	6.03×10^{-13}	1.28×10^{-12}	2.01×10^{-20}	5.14×10^{-20}	−2.71	Y
17	1.62×10^{-15}	3.07×10^{-15}	7.19×10^{-22}	1.03×10^{-21}	−3.03	Y
18	4.31×10^{-17}	9.02×10^{-17}	2.19×10^{-22}	5.72×10^{-22}	−2.75	Y
19	6.38×10^{-13}	1.80×10^{-12}	3.19×10^{-20}	6.17×10^{-20}	−2.03	Y
20	2.66×10^{-12}	6.49×10^{-12}	3.29×10^{-20}	5.28×10^{-20}	−2.36	Y
21	1.46×10^{-12}	7.82×10^{-12}	1.37×10^{-20}	4.70×10^{-20}	−1.07	N
22	1.47×10^{-12}	2.63×10^{-12}	4.76×10^{-20}	1.54×10^{-19}	−3.22	Y
23	1.30×10^{-12}	3.46×10^{-12}	1.82×10^{-20}	3.14×10^{-20}	−2.15	Y
24	6.63×10^{-13}	1.39×10^{-12}	3.25×10^{-20}	6.89×10^{-20}	−2.74	Y
25	1.65×10^{-15}	5.19×10^{-15}	9.69×10^{-22}	2.29×10^{-21}	−1.82	Y
26	2.40×10^{-8}	1.36×10^{-8}	6.05×10^{-18}	4.05×10^{-18}	−10.17	Y
27	1.57×10^{-13}	2.55×10^{-13}	7.03×10^{-20}	1.14×10^{-19}	−3.54	Y
28	3.95×10^{-13}	9.97×10^{-13}	5.67×10^{-20}	1.09×10^{-19}	−2.28	Y
29	2.18×10^{-12}	6.80×10^{-12}	4.76×10^{-20}	9.85×10^{-20}	−1.84	Y
30	1.56×10^{-12}	3.32×10^{-12}	3.38×10^{-20}	6.08×10^{-20}	−2.7	Y
31	1.61×10^{-12}	3.89×10^{-12}	6.30×10^{-20}	7.13×10^{-20}	−2.38	Y
32	1.37×10^{-12}	3.30×10^{-12}	2.71×10^{-20}	6.28×10^{-20}	−2.39	Y
33	2.32×10^{-12}	6.87×10^{-12}	4.24×10^{-20}	8.03×10^{-20}	−1.94	Y
34	1.09×10^{-20}	3.96×10^{-20}	4.68×10^{-24}	2.31×10^{-24}	−1.58	N
35	1.79×10^{-12}	5.13×10^{-12}	2.81×10^{-20}	4.92×10^{-20}	−2.01	Y
36	1.25×10^{-12}	3.25×10^{-12}	9.93×10^{-20}	2.71×10^{-19}	−2.22	Y

33

33

Table 11. Hypothesis test T1FLS vs. IT2FLS—50 dimensions.

No	T1FLS		IT2FLS		Hypothesis Test	
	fisGau318 50		it2_3gausS6523 50			
	Mean	SD	Mean	SD	Z	E
1	1.30×10^{-12}	3.49×10^{-12}	3.76×10^{-20}	9.84×10^{-20}	-2.04	Y
2	7.96×10^{-16}	2.01×10^{-15}	7.56×10^{-22}	9.99×10^{-22}	-2.17	Y
3	1.90×10^{-16}	5.43×10^{-16}	1.81×10^{-22}	4.36×10^{-22}	-1.92	Y
4	8.02×10^{-13}	1.60×10^{-12}	4.19×10^{-20}	8.97×10^{-20}	-2.74	Y
5	2.48×10^{-9}	1.35×10^{-8}	2.80×10^{-18}	5.28×10^{-18}	-1	N
6	1.47×10^{-12}	3.39×10^{-12}	2.80×10^{-20}	3.96×10^{-20}	-2.38	Y
7	1.46×10^{-12}	4.47×10^{-12}	5.97×10^{-20}	1.20×10^{-19}	-1.78	Y
8	1.45×10^{-12}	4.24×10^{-12}	7.59×10^{-20}	2.55×10^{-19}	-1.88	Y
9	1.01×10^{-12}	2.80×10^{-12}	4.39×10^{-20}	7.03×10^{-20}	-1.97	Y
10	9.81×10^{-13}	2.42×10^{-12}	4.86×10^{-20}	8.13×10^{-20}	-2.22	Y
11	4.46×10^{-13}	7.81×10^{-13}	4.14×10^{-20}	6.42×10^{-20}	-3.13	Y
12	6.76×10^{-15}	1.34×10^{-14}	1.64×10^{-20}	5.89×10^{-20}	-2.76	Y
13	1.48×10^{-12}	3.89×10^{-12}	5.64×10^{-20}	9.65×10^{-20}	-2.09	Y
14	1.99×10^{-12}	7.02×10^{-12}	2.63×10^{-20}	4.20×10^{-20}	-1.55	N
15	4.11×10^{-13}	9.33×10^{-13}	2.63×10^{-20}	4.71×10^{-20}	-2.41	Y

Table 11. Cont.

No	T1FLS		IT2FLS		Hypothesis Test	
	fisGau318 50		it2_3gausS6523 50			
	Mean	SD	Mean	SD	Z	E
16	5.79×10^{-13}	1.11×10^{-12}	2.11×10^{-20}	2.65×10^{-20}	-2.87	Y
17	5.12×10^{-16}	8.36×10^{-16}	1.06×10^{-21}	1.78×10^{-21}	-3.36	Y
18	1.46×10^{-16}	3.83×10^{-16}	1.78×10^{-22}	7.21×10^{-22}	-2.09	Y
19	4.32×10^{-12}	1.65×10^{-11}	2.90×10^{-20}	5.45×10^{-20}	-1.44	N
20	7.51×10^{-13}	1.86×10^{-12}	3.47×10^{-20}	5.03×10^{-20}	-2.21	Y
21	4.73×10^{-13}	2.32×10^{-12}	6.02×10^{-21}	1.12×10^{-20}	-1.12	N
22	8.50×10^{-13}	3.52×10^{-12}	6.65×10^{-20}	1.92×10^{-19}	-1.32	N
23	5.72×10^{-13}	1.03×10^{-12}	1.05×10^{-19}	1.74×10^{-19}	-3.06	Y
24	3.60×10^{-12}	1.83×10^{-11}	3.34×10^{-20}	6.81×10^{-20}	-1.08	N
25	1.43×10^{-15}	4.36×10^{-15}	9.72×10^{-22}	2.36×10^{-21}	-1.8	Y
26	1.79×10^{-8}	1.05×10^{-8}	7.26×10^{-18}	5.36×10^{-18}	-9.35	Y
27	3.87×10^{-13}	8.95×10^{-13}	4.36×10^{-20}	8.94×10^{-20}	-2.37	Y
28	4.16×10^{-13}	8.78×10^{-13}	4.68×10^{-20}	1.01×10^{-19}	-2.6	Y
29	2.97×10^{-13}	6.44×10^{-13}	3.95×10^{-20}	6.92×10^{-20}	-2.53	Y
30	4.08×10^{-13}	7.22×10^{-13}	3.06×10^{-20}	6.69×10^{-20}	-3.09	Y
31	2.99×10^{-12}	1.13×10^{-11}	5.60×10^{-20}	1.23×10^{-19}	-1.46	N
32	1.67×10^{-13}	4.69×10^{-13}	4.81×10^{-20}	9.99×10^{-20}	-1.95	Y
33	1.76×10^{-12}	4.06×10^{-12}	2.74×10^{-20}	4.33×10^{-20}	-2.37	Y
34	3.74×10^{-20}	1.16×10^{-19}	6.08×10^{-24}	3.46×10^{-24}	-1.76	Y
35	3.83×10^{-13}	7.54×10^{-13}	2.64×10^{-20}	3.01×10^{-20}	-2.78	Y
36	6.11×10^{-13}	1.02×10^{-12}	6.27×10^{-20}	1.32×10^{-19}	-3.27	Y

29

Table 12. Hypothesis test T1FLS vs. IT2FLS—100 dimensions.

No	T1FLS		IT2FLS		Hypothesis Test	
	fisGau318 100		it2_3gausS6523 100			
	Mean	SD	Mean	SD	Z	E
1	3.86×10^{-12}	1.66×10^{-11}	3.97×10^{-20}	6.97×10^{-20}	-1.27	N
2	7.39×10^{-15}	2.63×10^{-14}	7.99×10^{-22}	1.81×10^{-21}	-1.54	N
3	3.01×10^{-14}	8.31×10^{-14}	6.00×10^{-21}	1.10×10^{-20}	-1.98	Y
4	8.71×10^{-13}	3.38×10^{-12}	6.70×10^{-20}	1.42×10^{-19}	-1.41	N
5	1.93×10^{-8}	1.30×10^{-8}	6.42×10^{-18}	4.23×10^{-18}	-8.11	Y
6	7.90×10^{-13}	2.03×10^{-12}	5.60×10^{-20}	7.87×10^{-20}	-2.13	Y
7	2.13×10^{-12}	6.74×10^{-12}	6.67×10^{-20}	2.76×10^{-19}	-1.73	Y
8	8.91×10^{-13}	2.71×10^{-12}	5.84×10^{-20}	1.36×10^{-19}	-1.8	Y
9	5.46×10^{-13}	1.17×10^{-12}	3.56×10^{-20}	6.70×10^{-20}	-2.55	Y
10	9.03×10^{-13}	1.97×10^{-12}	3.56×10^{-20}	6.07×10^{-20}	-2.51	Y
11	1.22×10^{-12}	3.91×10^{-12}	6.20×10^{-20}	8.90×10^{-20}	-1.71	Y
12	2.71×10^{-20}	7.10×10^{-20}	6.23×10^{-24}	4.84×10^{-24}	-2.09	Y
13	2.11×10^{-12}	5.48×10^{-12}	1.70×10^{-20}	2.93×10^{-20}	-2.11	Y
14	4.24×10^{-13}	8.07×10^{-13}	3.15×10^{-20}	6.48×10^{-20}	-2.88	Y
15	1.28×10^{-13}	2.52×10^{-13}	4.77×10^{-20}	8.87×10^{-20}	-2.78	Y
16	5.58×10^{-13}	1.19×10^{-12}	2.86×10^{-20}	3.96×10^{-20}	-2.57	Y
17	4.69×10^{-16}	1.66×10^{-15}	8.02×10^{-22}	1.53×10^{-21}	-1.54	N
18	8.38×10^{-17}	2.94×10^{-16}	5.96×10^{-23}	1.08×10^{-22}	-1.56	N
19	9.12×10^{-13}	1.68×10^{-12}	3.73×10^{-20}	4.39×10^{-20}	-2.98	Y
20	7.58×10^{-13}	1.62×10^{-12}	2.33×10^{-20}	3.70×10^{-20}	-2.57	Y
21	2.98×10^{-14}	6.60×10^{-14}	1.78×10^{-20}	4.97×10^{-20}	-2.47	Y
22	1.09×10^{-12}	3.49×10^{-12}	1.49×10^{-20}	2.85×10^{-20}	-1.71	Y
23	3.76×10^{-13}	8.72×10^{-13}	1.84×10^{-20}	3.05×10^{-20}	-2.36	Y
24	2.26×10^{-12}	7.21×10^{-12}	2.24×10^{-20}	4.26×10^{-20}	-1.71	Y
25	6.01×10^{-16}	1.50×10^{-15}	3.30×10^{-22}	6.23×10^{-22}	-2.19	Y
26	1.87×10^{-8}	9.88×10^{-9}	7.70×10^{-18}	5.55×10^{-18}	-10.36	Y
27	2.11×10^{-12}	7.50×10^{-12}	1.53×10^{-20}	2.38×10^{-20}	-1.54	N
28	4.11×10^{-12}	9.48×10^{-12}	4.93×10^{-20}	8.81×10^{-20}	-2.37	Y
29	6.45×10^{-13}	1.02×10^{-12}	3.25×10^{-20}	5.28×10^{-20}	-3.45	Y
30	2.99×10^{-13}	5.03×10^{-13}	2.19×10^{-20}	4.23×10^{-20}	-3.25	Y

Table 12. Cont.

No	T1FLS		IT2FLS		Hypothesis Test	
	fisGau318 100		it2_3gausS6523 100			
	Mean	SD	Mean	SD	Z	E
31	9.34×10^{-13}	2.65×10^{-12}	2.12×10^{-20}	3.97×10^{-20}	−1.93	Y
32	1.71×10^{-13}	5.44×10^{-13}	2.56×10^{-20}	4.14×10^{-20}	−1.72	Y
33	1.33×10^{-12}	3.02×10^{-12}	2.55×10^{-20}	3.25×10^{-20}	−2.41	Y
34	5.56×10^{-20}	2.29×10^{-19}	4.96×10^{-24}	2.05×10^{-24}	−1.33	N
35	4.14×10^{-13}	8.86×10^{-13}	1.23×10^{-19}	3.19×10^{-19}	−2.56	Y
36	1.55×10^{-12}	4.83×10^{-12}	2.53×10^{-20}	4.93×10^{-20}	−1.76	Y
						29

Figures 27 and 28 show the behaviors of the mean and standard deviation of the two types of fuzzy systems T1FLS vs. IT2FLS, for 30 dimensions.

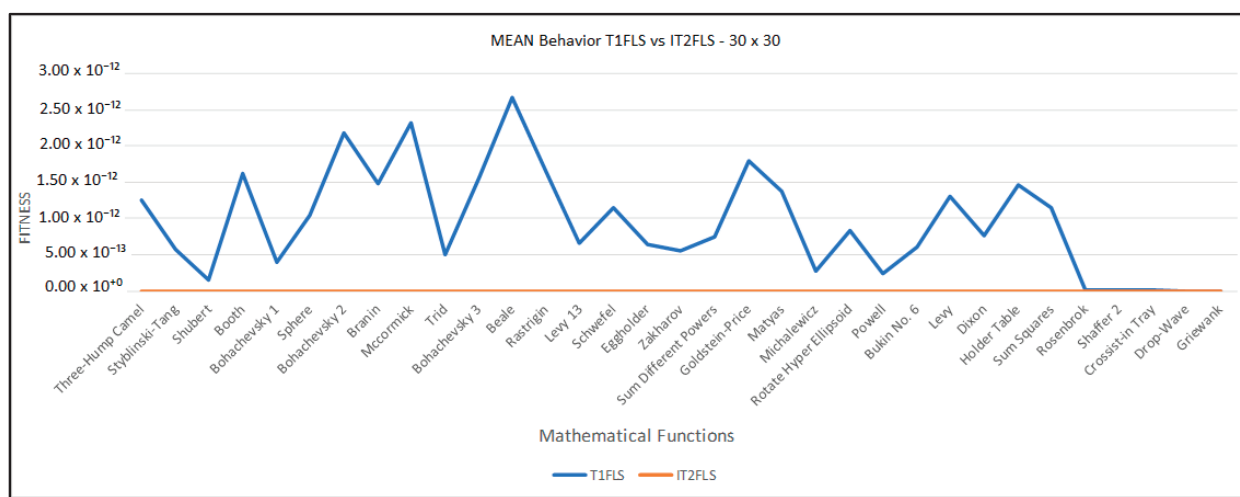


Figure 27. Behavior of the MEAN for T1FLS vs. IT2FLS 30 dimensions.

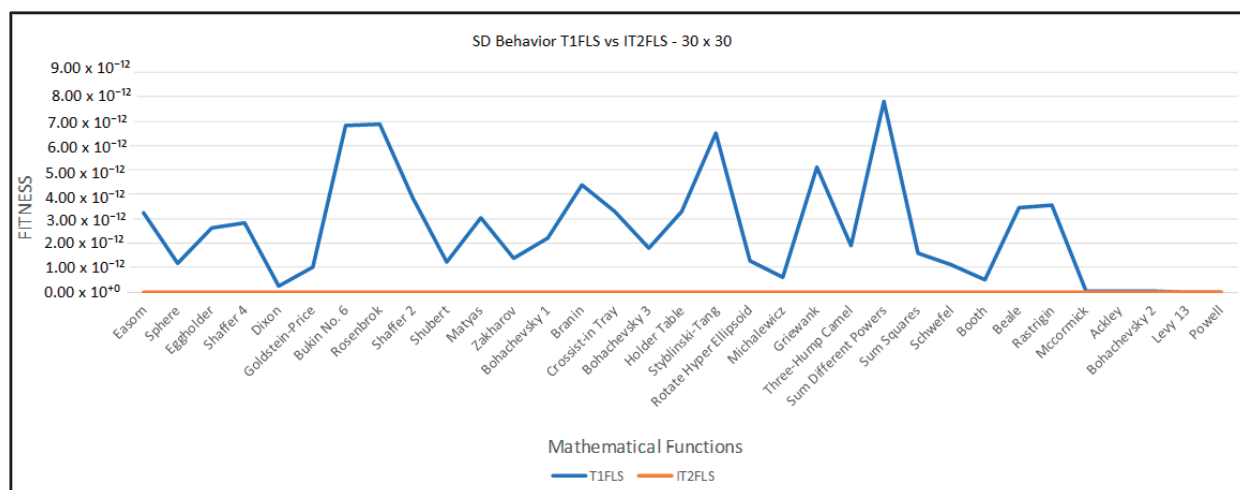


Figure 28. Behavior of the SD for T1FLS vs. IT2FLS 30 dimensions.

Comparison Table 11 shows the hypothesis test for 50 dimensions of the fuzzy systems T1FLS with parameter adaptation and optimized IT2FLS was better in 29 out of 36 mathematical functions.

Figures 29 and 30 show the behaviors of the mean and standard deviation of the two types of fuzzy systems T1FLS vs. IT2FLS, for 50 dimensions and 30 iterations.

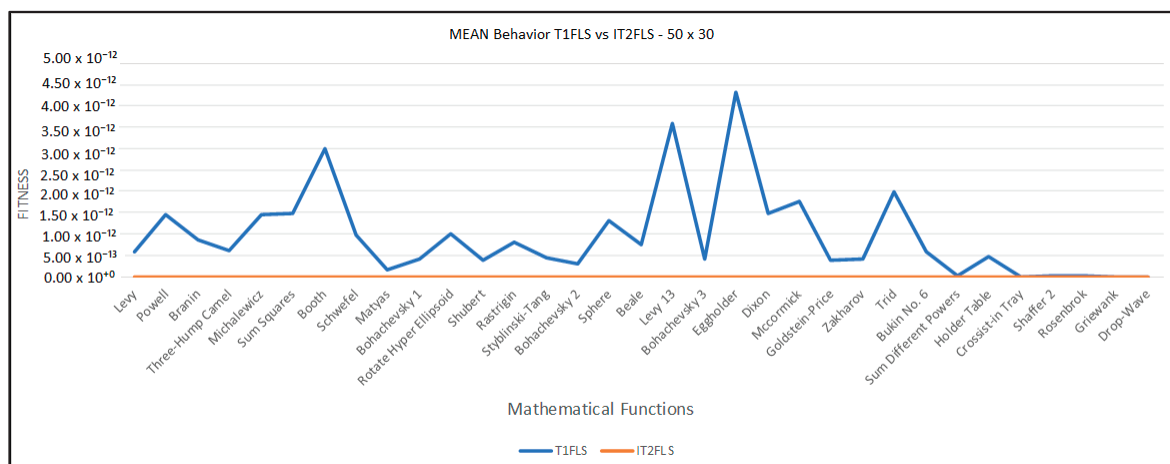


Figure 29. Behavior of the MEAN for T1FLS vs. IT2FLS 50 dimensions.

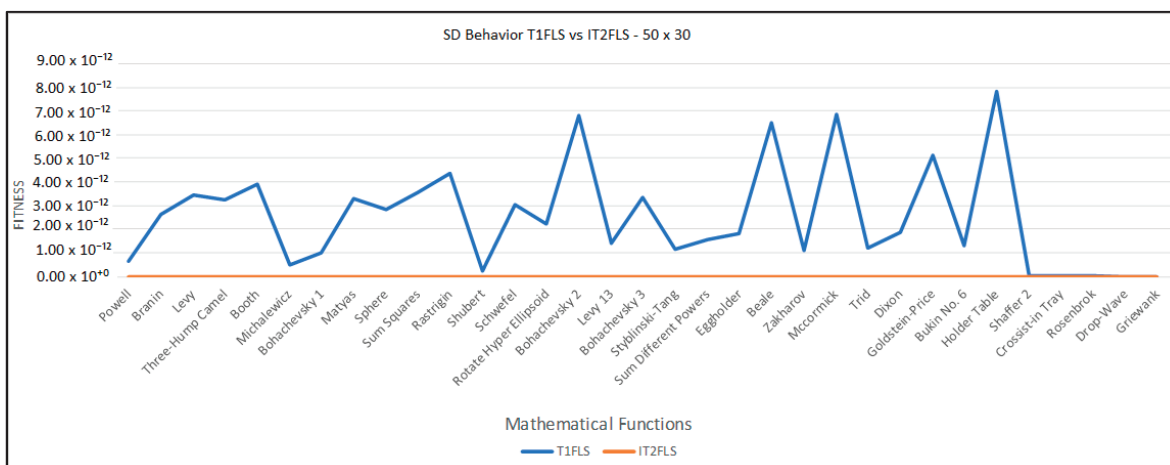


Figure 30. Behavior of the SD for T1FLS vs. IT2FLS 50 dimensions.

Comparison Table 12 shows the hypothesis test for 100 dimensions of the fuzzy systems T1FLS with parameter adaptation and optimized IT2FLS was better in 29 out of 36 mathematical functions.

Figures 31 and 32 show the behaviors of the mean and standard deviation of the two types of fuzzy systems T1FLS vs. IT2FLS for 100 dimensions.

6.2. Discussion of Results

The use of metaheuristics in model optimization is a constant in all research works in artificial intelligence. In this research, we used the DMOA optimization algorithm, T1FLS and IT2FLS fuzzy systems with parameter adaptation, 36 mathematical functions, from CEC-2013 (Table 2) to measure their performance capabilities, and hypothesis testing was performed to test the optimization capability.

When evaluating the results of the hypothesis tests T1FLS and IT2FLS fuzzy systems for 30 dimensions, both with parameter adaptation in Table 7, the results favor the T1FLS fuzzy system in 31 out of 36 mathematical functions evaluated, however when optimizing the parameters of the membership functions of the fuzzy IT2FLS system for 30 dimensions the results of the hypothesis tests favor IT2FLS in 33 out of 36 hypothesis tests in Table 10, and for 50 dimensions the hypothesis tests favor IT2FLS in 29 out of 36 hypothesis tests in Table 11 and finally for 100 dimensions the results favor IT2FLS in 29 out of 36 hypothesis tests in Table 12.

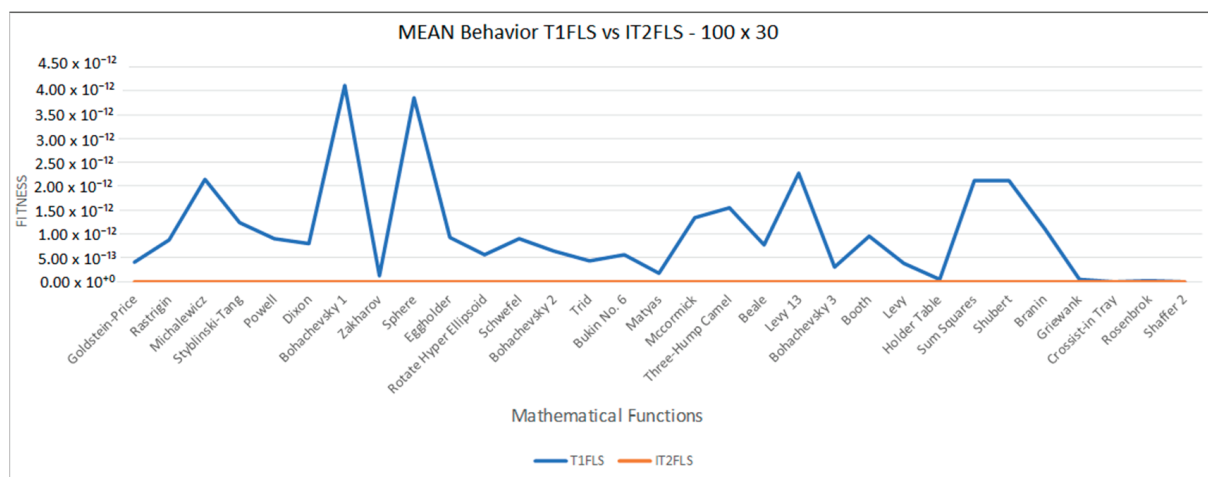


Figure 31. Behavior of the MEAN for T1FLS vs. IT2FLS 100 dimensions.

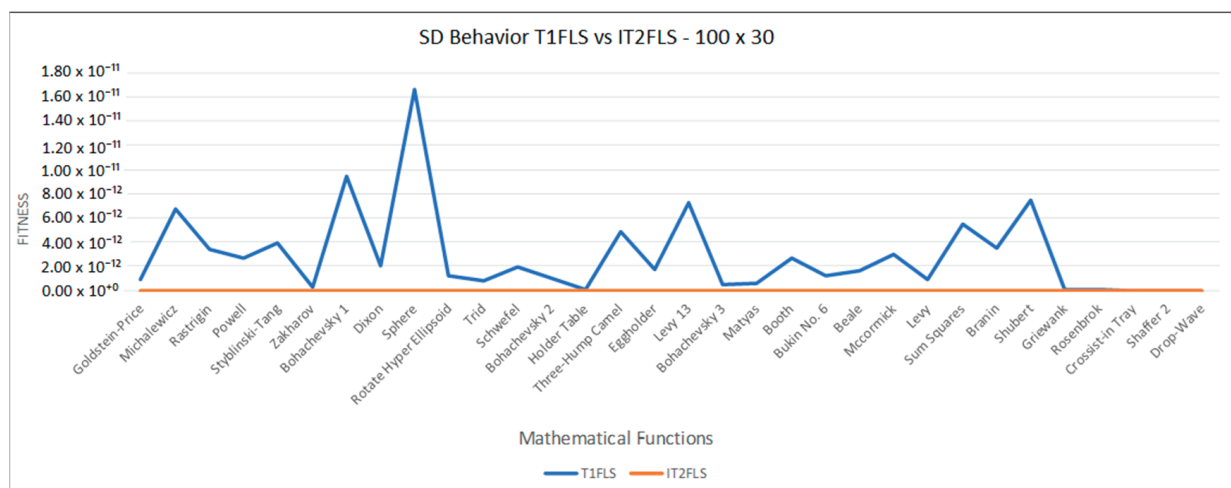


Figure 32. Behavior of the SD for T1FLS vs. IT2FLS 100 dimensions.

6.3. Programming Environment

The language used in the programming of the DMOA algorithm is MATLAB R2017b and the equipment where the programming and experiments were carried out is a Desktop Computer Intel Core i5 4460S 2.90 GHz., RAM memory DDR3 16 Gb, Intel HD Graphics 4600, and Operating System Windows 10 Professional.

7. Conclusions

As we have already mentioned, the DMOA is a metaheuristic that simulates a biological system that exists in nature, where there is diversity of plants, and where there is communication among plants and the Mycorrhiza Network. To model this ecosystem, the discrete Lotka–Volterra models can be used: the defense model that simulates emergency situations that can manifest in the ecosystem, the cooperative model that simulates the exchange of resources such as CO₂, water, nitrogen, phosphorus, potassium, etc., and the competitive model that represents how the Mycorrhizal Network can be extended by adding other plants competing in the habitat for resources and how within the network larger plants and the Mycorrhiza Network offer resources to growing plants.

Experiments and statistical tests were carried out with the DMOA optimization algorithm with 36 CEC-2013 mathematical functions, as can be seen in Table 2. The configuration of the parameters significantly affects the performance of the algorithm and therefore its convergence and it is very important to find the relationship between parameter values and

convergence rate and adjust them for better performance. The statistical tests carried out tell us that when we optimize the parameters of the membership functions of the IT2FLS fuzzy systems with the DMOA algorithm, we obtain better results than the T1FLS fuzzy systems with parameter adaptation: in 33 of 36 statistical tests for 30 dimensions, shown in Table 10, in 29 of 36 statistical tests for 50 dimensions (shown in Table 11) and in 29 of 36 tests for 100 dimensions (shown in Table 12). In summary, of 108 statistical tests carried out in 91 tests, the IT2FLS fuzzy systems optimized with the DMOA algorithm are better, which is in 84.25% of the cases.

We have previously applied the DMOA algorithm in the optimization of the architecture of a non-linear autoregressive neural network for Mackey–Glass time series prediction [69], and in this article in the adaptation of the parameters of the fuzzy systems T1FLS and IT2FLS. Additionally, hypothesis tests were carried out and the results obtained in both investigations were favorable for the DMOA optimization algorithm. In the near future, we plan to conduct research applying the CMOA (Continuous Mycorrhiza Optimization Algorithm) and DMOA in the optimization of the architecture of a long short-term memory (LSTM) neural network. We also plan to solve control problems with the two algorithms (CMOA-DMOA) and systems: type-1 fuzzy logic system (T1FLS), interval type-2 fuzzy logic system (IT2FLS) and generalized type-2 fuzzy logic system (GT2FLS). In addition, experiments with the mathematical functions of CEC-2017 and CEC-2019 are also planned. These two optimization methods, CMOA and DMOA, can be very useful in the optimization of neural networks and fuzzy systems architectures in an efficient way due to their fast convergence, as we have already seen with NARNN neural networks and IT2FLS fuzzy systems, and for this reason we expect to obtain good results in the future research works that we intend to perform.

Author Contributions: Conceptualization, F.V. and O.C.; methodology, F.V.; software, H.C.-O.; validation, H.C.-O. and F.V.; formal analysis, H.C.-O.; investigation, O.C.; resources, O.C.; writing—original draft preparation, F.V. and O.C.; writing—review and editing, H.C.-O.; visualization, F.V.; supervision, F.V. and O.C.; project administration. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Data Availability Statement: Not applicable.

Conflicts of Interest: All the authors in the paper have no conflict of interest.

References

1. Yang, X.-S. Cuckoo Search and Firefly Algorithm: Overview and Analysis. In *Studies in Computational Intelligence*; Springer: Cham, Switzerland; Berlin/Heidelberg, Germany; New York, NY, USA; Dordrecht, The Netherlands; London, UK, 2014; pp. 1–26.
2. Yang, X.-S.; Deb, S. Cuckoo Search via Lévy flights. In *Proceedings of the 2009 World Congress on Nature & Biologically Inspired Computing (NaBIC)*, Coimbatore, India, 9–11 December 2009; pp. 210–214.
3. Rosselan, M.; Sulaiman, S.I.; Othman, N. Evaluation of Fast Evolutionary Programming, Firefly Algorithm and Mutate-Cuckoo Search Algorithm in Single-Objective Optimization. *Int. J. Electr. Electron. Syst. Res.* **2019**, *9*, 1–6. [CrossRef]
4. Yang, X.-S.; Deb, S.; Mishra, S.K. Multi-species Cuckoo Search Algorithm for Global Optimization. *Cogn. Comput.* **2018**, *10*, 1085–1095. [CrossRef]
5. Yang, X.-S.; Deb, S. Cuckoo search: Recent advances and applications. *Neural Comput. Appl.* **2014**, *24*, 169–174. [CrossRef]
6. Yang, X.-S.; Koziel, S. Computational Optimization: An Overview. In *Computational Optimization, Methods and Algorithms*; Koziel, S., Yang, X.-S., Eds.; Springer: Berlin/Heidelberg, Germany, 2011; pp. 1–11.
7. Yang, X.-S. *Engineering Optimization: An Introduction with Metaheuristic Applications*; University of Cambridge, Department of Engineering: Cambridge, UK; John Wiley & Sons, Inc.: Hoboken, NJ, USA, 2010; pp. 15–27.
8. Abdel-Basset, M.; Shawky, L.; Kumar, A. A comparative study of cuckoo search and flower pollination algorithm on solving global optimization problems. *Libr. Hi Tech.* **2017**, *35*, 588–601. [CrossRef]
9. El-Ela, A.A.A.; Mouwafi, M.T.; Elbaset, A.A. *Modern Optimization Techniques for Smart Grids*, Department of Electromechanics Engineering, 1st ed.; Springer International Publishing: Cham, Switzerland, 2023.
10. Tyagi, M.; Sachdeva, A.; Sharma, V. *Optimization Methods in Engineering, Lecture Notes on Multidisciplinary Industrial Engineering*; Springer: Singapore, 2021.
11. Yakimov, A.S. *Analytical Solution Methods for Boundary Value Problems*, 1st ed.; Nikki Levy: London, UK, 2016.

12. Badar, A.Q.H. *Evolutionary Optimization Algorithms*, 1st ed.; CRC Press: Boca Raton, FL, USA, 2021.
13. Qu, C.; He, W.; Peng, X.; Peng, X. Harris Hawks optimization with information exchange. *Appl. Math. Model.* **2020**, *84*, 52–75. [CrossRef]
14. Mao, K.; Pan, Q.; Pang, X.; Chai, T. A novel Lagrangian relaxation approach for a hybrid flowshop scheduling problem in the steelmaking-continuous casting process. *Eur. J. Oper. Res.* **2014**, *236*, 51–60. [CrossRef]
15. Poli, R.; Kennedy, J.; Blackwell, T. Particle swarm optimization. *Swarm Intell.* **2007**, *1*, 33–57. [CrossRef]
16. Gandomi, A.H.; Yang, X.-S.; Alavi, A.H.; Talatahari, S. Bat algorithm for constrained optimization tasks. *Neural Comput. Appl.* **2013**, *22*, 1239–1255. [CrossRef]
17. Jingqiao, Z.; Sanderson, A.C. JADE: Adaptive Differential Evolution with Optional External Archive. *IEEE Trans. Evol. Comput.* **2009**, *13*, 945–958. [CrossRef]
18. Rajabioun, R. Cuckoo Optimization Algorithm. *Appl. Soft Comput.* **2011**, *11*, 5508–5518. [CrossRef]
19. Yang, X.-S. *Flower Pollination Algorithm for Global Optimization*; Durand-Lose, J., Jonoska, N., Eds.; Unconventional Computation and Natural Computation, UCNC 2012; Lecture Notes in Computer Science; Springer: Berlin/Heidelberg, Germany, 2012; Volume 7445, pp. 240–249. [CrossRef]
20. Askarzadeh, A. A novel metaheuristic method for solving constrained engineering optimization problems: Crow search algorithm. *Comput. Struct.* **2016**, *169*, 1–12. [CrossRef]
21. Mirjalili, S.; Mirjalili, S.M.; Lewis, A. Grey Wolf Optimizer. *Adv. Eng. Softw.* **2014**, *69*, 46–61. [CrossRef]
22. Rao, R.V.; Savsani, V.J.; Vakharia, D.P. Teaching–Learning–Based Optimization: An optimization method for continuous non-linear large scale problems. *Inf. Sci. (N. Y.)* **2012**, *183*, 1–15. [CrossRef]
23. Cheng, M.-Y.; Prayogo, D. Symbiotic Organisms Search: A new metaheuristic optimization algorithm. *Comput. Struct.* **2014**, *139*, 98–112. [CrossRef]
24. Moayedi, H.; Mu’azu, M.A.; Foong, L.K. Novel swarm-based approach for predicting the cooling load of residential buildings based on social behavior of elephant herds. *Energy Build.* **2020**, *206*, 109579. [CrossRef]
25. Benaissa, B.; Hocine, N.A.; Khatir, S.; Riahi, M.K.; Mirjalili, S. YUKI Algorithm and POD-RBF for Elastostatic and dynamic crack identification. *J. Comput. Sci.* **2021**, *55*, 101451. [CrossRef]
26. Khatir, A.; Capozucca, R.; Khatir, S.; Magagnini, E.; Benaissa, B.; Le Thanh, C.; Abdel Wahab, M. A new hybrid PSO-YUKI for double cracks identification in CFRP cantilever beam. *Compos. Struct.* **2023**, *311*, 116803. [CrossRef]
27. Amoura, N.; Benaissa, B.; Al Ali, M.; Khatir, S. *Deep Neural Network and YUKI Algorithm for Inner Damage Characterization Based on Elastic Boundary Displacement*; Capozucca, R., Khatir, S., Milani, G., Eds.; International Conference of Steel and Composite for Engineering Structures, ICSCS 2022; Lecture Notes in Civil Engineering; Springer: Cham, Switzerland, 2023; Volume 317, pp. 220–233. [CrossRef]
28. Hashim, F.A.; Hussien, A.G. Snake Optimizer: A novel meta-heuristic optimization algorithm. *Knowl. Based Syst.* **2022**, *242*, 108320. [CrossRef]
29. Diwekar, U.M. *Introduction to Applied Optimization*, 3rd ed.; Springer International Publishing: Cham, Switzerland, 2020.
30. Ghaemi, M.B.; Gharakhanlu, N.; Rassias, T.M.; Saadati, R. *Advances in Matrix Inequalities*, 1st ed.; Springer International Publishing: Cham, Switzerland, 2021.
31. Lange, K. *Optimization*, 2nd ed.; Springer: New York, NY, USA, 2013.
32. Kochenderfer, M.J.; Wheeler, T.A. *Algorithms for Optimization*, 1st ed.; The MIT Press: Cambridge, MA, USA, 2019.
33. Adam, S.P.; Alexandropoulos, S.-A.N.; Pardalos, P.M.; Vrahatis, M.N. No Free Lunch Theorem: A Review. In *Approximation and Optimization: Algorithms, Complexity and Applications*; Demetriou, I.C., Pardalos, P.M., Eds.; Springer International Publishing: Cham, Switzerland, 2019; pp. 57–82.
34. Zadeh, L.A. Fuzzy Sets. *Inf. Control.* **1965**, *8*, 338–353. [CrossRef]
35. Zadeh, L.A.; Fu, K.S.; Tanaka, K. *Fuzzy Sets and their Applications to Cognitive and Decision Processes*; Academic Press: London, UK, 1975; pp. 1–39.
36. Zadeh, L.A. Toward a theory of fuzzy information granulation and its centrality in human reasoning and fuzzy logic. *Fuzzy Sets Syst.* **1997**, *90*, 111–127. [CrossRef]
37. Zadeh, L.A. The concept of a linguistic variable and its application to approximate reasoning—I. *Inf. Sci. (N. Y.)* **1975**, *8*, 199–249. [CrossRef]
38. Dubois, D.P.H. *Fuzzy Sets and Systems: Theory and Applications*, 1st ed.; Academic Press: New York, NY, USA, 1980.
39. Carreon, H.; Valdez, F. Comparative Study of Type-1 and Interval Type-2 Fuzzy Systems in Parameter Adaptation of the Fuzzy Flower Pollination Algorithm. In *Recent Advances of Hybrid Intelligent Systems Based on Soft Computing*, 1st ed.; Melin, P., Castillo, O., Kacprzyk, J., Eds.; Studies in Computational Intelligence; Springer: Cham, Switzerland, 2021; Volume 915, pp. 145–161. [CrossRef]
40. Hisdal, E. The IF THEN ELSE statement and interval-valued fuzzy sets of higher type. *Int. J. Man. Mach. Stud.* **1981**, *15*, 385–455. [CrossRef]
41. Liang, Q.; Mendel, J.M. Interval type-2 fuzzy logic systems: Theory and design. *IEEE Trans. Fuzzy Syst.* **2000**, *8*, 535–550. [CrossRef] [PubMed]

42. Karnik, N.N.; Mendel, J.M. Introduction to type-2 fuzzy logic systems. In Proceedings of the 1998 IEEE International Conference on Fuzzy Systems Proceedings, IEEE World Congress on Computational Intelligence (Cat. No.98CH36228), Anchorage, AK, USA, 4–9 May 1998; pp. 915–920.
43. Schueffler, A.; Anke, T. Fungal natural products in research and development. *Nat. Prod. Rep.* **2014**, *31*, 1425–1448. [CrossRef]
44. Smith, S.; Read, D. *Mycorrhizal Symbiosis*, 3rd ed.; Elsevier: New York, NY, USA, 2008.
45. Prasad, R.; Bhola, D.; Akdi, K.; Cruz, C.; KVSS, S.; Tuteja, N.; Varma, A. Introduction to Mycorrhiza: Historical Development. In *Mycorrhiza—Function, Diversity, State of the Art*, 4th ed.; Springer International Publishing: Cham, Switzerland, 2017; pp. 1–7.
46. Buscot, F. Implication of evolution and diversity in arbuscular and ectomycorrhizal symbioses. *J. Plant Physiol.* **2015**, *172*, 55–61. [CrossRef]
47. Zhao, M.; Xuan, Z.; Li, C. Dynamics of a discrete-time predator-prey system. *Adv. Differ. Equ.* **2016**, *2016*, 191. [CrossRef]
48. Saha, P.; Bairagi, N.; Biswas, M. On the Dynamics of a Discrete Predator–Prey Model. In *Trends in Biomathematics: Modeling, Optimization and Computational Problems*; Springer International Publishing: Cham, Switzerland, 2018; pp. 219–232.
49. Liu, P.; Elaydi. Discrete Competitive and Cooperative Models of Lotka–Volterra Type. *J. Comput. Anal. Appl.* **2001**, *3*, 53–73. [CrossRef]
50. Din, Q. Dynamics of a discrete Lotka–Volterra model. *Adv. Differ. Equ.* **2013**, *2013*, 95. [CrossRef]
51. El-Shorbagy, M.A.; Hassanien, A.E. Particle Swarm Optimization from Theory to Applications. *Int. J. Rough Sets Data Anal.* **2018**, *5*, 1–24. [CrossRef]
52. Balasubramani, K.; Marcus, K. A Study on Flower Pollination Algorithm and Its Applications. *Int. J. Appl. Innov. Eng. Manag. (IJAIEM)* **2014**, *3*, 230–235.
53. Dorigo, M.; Birattari, M.; Stutzle, T. Ant colony optimization. *IEEE Comput. Intell. Mag.* **2006**, *1*, 28–39. [CrossRef]
54. Bansal, J.C.; Sharma, H.; Jadon, S.S. Artificial bee colony algorithm: A survey. *Int. J. Adv. Intell. Paradig.* **2013**, *5*, 123. [CrossRef]
55. Ali, N.; Othman, M.A.; Husain, M.N.; Misran, M. A review of firefly algorithm. *ARPN J. Eng. Appl. Sci.* **2014**, *9*, 1732–1736.
56. Simard, S.W. Mycorrhizal Networks Facilitate Tree Communication, Learning, and Memory. In *Memory and Learning in Plants; Signaling and Communication in Plants*; Baluska, F., Gagliano, M., Witzany, G., Eds.; Springer: Cham, Switzerland, 2018.
57. Castro-Delgado, A.L.; Elizondo-Mesén, S.; Valladares-Cruz, Y.; Rivera-Méndez, W. Wood Wide Web: Communication through the mycorrhizal network. *Tecnol. Eng. Marcha J.* **2020**, *33*, 114–125.
58. Beiler, K.J.; Simard, S.W.; Durall, D.M. Topology of tree-mycorrhizal fungus interaction networks in xeric and mesic Douglas-fir forests. *J. Ecol.* **2015**, *103*, 616–628. [CrossRef]
59. Simard, S.W.; Asay, A.; Beiler, K.; Bingham, M.; Deslippe, J.; He, X.; Philip, L.; Song, Y.; Teste, F. Resource Transfer between Plants Through Ectomycorrhizal Fungal Networks. In *Mycorrhizal Networks*; Horton, T., Ed.; Ecological Studies; Springer: Dordrecht, The Netherlands, 2015; p. 224.
60. Gorzelak, M.A.; Asay, A.K.; Pickles, B.J.; Simard, S.W. Inter-plant communication through mycorrhizal networks mediates complex adaptive behaviour in plant communities. *AoB PLANTS* **2015**, *7*, plv050. [CrossRef]
61. Mehrabian, A.R.; Lucas, C. A novel numerical optimization algorithm inspired from weed colonization. *Ecol. Inform.* **2006**, *1*, 355–366. [CrossRef]
62. Premaratne, U.; Samarabandu, J.; Sidhu, T. A new biologically inspired optimization algorithm. In Proceedings of the 2009 International Conference on Industrial and Information Systems (ICIIS), Peradeniya, Sri Lanka, 28–31 December 2009; pp. 279–284.
63. Salhi, A.; Fraga, E. *Nature-Inspired Optimisation Approaches and the New Plant Propagation Algorithm*; Computer Science: London, UK, 2011.
64. Novoplansky, A. Developmental plasticity in plants: Implications of non-cognitive behavior. *Evol. Ecol.* **2002**, *16*, 177–188. [CrossRef]
65. Akyol, S.; Alatas, B. Plant intelligence based metaheuristic optimization algorithms. *Artif. Intell. Rev.* **2017**, *47*, 417–462. [CrossRef]
66. Caraveo, C.; Valdez, F.; Castillo, O. *A New Bio-Inspired Optimization Algorithm Based on the Self-Defense Mechanisms of Plants*; Melin, P., Castillo, O., Kacprzyk, J., Eds.; Design of Intelligent Systems Based on Fuzzy Logic, Neural Networks and Nature-Inspired Optimization, Studies in Computational Intelligence; Springer: Cham, Switzerland, 2015; Volume 601, pp. 211–218. [CrossRef]
67. Carreon, H.; Valdez, F.; Castillo, O. Fuzzy Flower Pollination Algorithm to Solve Control Problems. In *Hybrid Intelligent Systems in Control, Pattern Recognition and Medicine*, 1st ed.; Castillo, O., Melin, P., Eds.; Springer: Cham, Switzerland, 2020; pp. 119–154.
68. Cai, W.; Yang, W.; Chen, X. A Global Optimization Algorithm Based on Plant Growth Theory: Plant Growth Optimization. In Proceedings of the 2008 International Conference on Intelligent Computation Technology and Automation (ICICTA), Changsha, China, 20–22 October 2008; pp. 1194–1199.
69. Carreon-Ortiz, H.; Valdez, F.; Melin, P.; Castillo, O. Architecture Optimization of a Non-Linear Autoregressive Neural Networks for Mackey–Glass Time Series Prediction Using Discrete Mycorrhiza Optimization Algorithm. *Micromachines* **2023**, *14*, 149. [CrossRef]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

Article

Multi-Objective Fault-Coverage Based Regression Test Selection and Prioritization Using Enhanced ACO_TCSP

Shweta Singhal ¹, Nishtha Jatana ^{2,*}, Kavita Sheoran ², Geetika Dhand ², Shaily Malik ², Reena Gupta ³, Bharti Suri ³, Mudligiriappa Niranjanamurthy ⁴, Sachi Nandan Mohanty ⁵ and Nihar Ranjan Pradhan ^{5,*}

¹ Department of Computer Science and Information Technology, Indira Gandhi Delhi Technical University for Women, New Delhi 110006, India; miss.shweta.singhal@gmail.com

² Department of Computer Science and Engineering, Maharaja Surajmal Institute of Technology, New Delhi 110058, India

³ University School of Information & Communication Technology, Guru Gobind Singh Indraprastha University, New Delhi 110078, India

⁴ Department of AI and ML, BMS Institute of Technology and Management, Bengaluru 560064, India

⁵ School of Computer Science & Engineering, VIT-AP University, Amaravati 522237, India; sachinandan09@gmail.com

* Correspondence: nishtha.jatana@gmail.com (N.J.); nihaar.pradhan@vitap.ac.in (N.R.P.)

Abstract: Regression testing of the software during its maintenance phase, requires test case prioritization and selection due to the dearth of the allotted time. The resources and the time in this phase are very limited, thus testers tend to use regression testing methods such as test case prioritization and selection. The current study evaluates the effectiveness of testing with two major goals: (1) Least running time and (2) Maximum fault coverage possible. Ant Colony Optimization (ACO) is a well-known soft computing technique that draws its inspiration from nature and has been widely researched, implemented, analyzed, and validated for regression test prioritization and selection. Many versions of ACO approaches have been prolifically applied to find solutions to many non-polynomial time-solvable problems. Hence, an attempt has been made to enhance the performance of the existing ACO_TCSP algorithm without affecting its time complexity. There have been efforts to enhance the exploration space of various paths in each iteration and with elite exploitation, reducing the total number of iterations required to converge to an optimal path. Counterbalancing enhanced exploration with intelligent exploitation implies that the run time is not adversely affected, the same has also been empirically validated. The enhanced algorithm has been compared with the existing ACO algorithm and with the traditional approaches. The approach has also been validated on four benchmark programs to empirically evaluate the proposed Enhanced ACO_TCSP algorithm. The experiment revealed the increased cost-effectiveness and correctness of the algorithm. The same has also been validated using the statistical test (independent *t*-test). The results obtained by evaluating the proposed approach against other reference techniques using Average Percentage of Faults Detected (APFD) metrics indicate a near-optimal solution. The multiple objectives of the highest fault coverage and least running time were fruitfully attained using the Enhanced ACO_TCSP approach without compromising the complexity of the algorithm.

Keywords: ant colony optimization; regression testing; test suite prioritization; metaheuristic technique; nature-inspired technique

MSC: 68-04

1. Introduction

Software has unprecedentedly altered the lifestyles of people in current and forthcoming times. In order to prevent software from getting superseded, lots of effort is required for its maintenance, thus incurring a hefty amount of money. Errors in software

can be functional or non-functional. Functional errors are related to the operations and actions of software, whereas non-functional errors are related to customer expectations and performance requirements. Function testing includes unit testing, system testing, and acceptance testing, whereas non-function testing deals with evaluating the parameters such as security, reliability, usability, and scalability. Software testing needed to ensure the proper functioning of the software after updates are referred to as Regression Testing [1]. Researchers have been rigorously working toward the development and validation of various regression testing techniques. Regression Test Prioritization and Regression Test Selection are the key regression testing techniques focused on in this paper. Regression test selection and prioritization activities, when carried out in an endeavor to obtain promising results in the minimum possible time, are an NP—complete combinatorial optimization problem [2]. The taxonomy of Regression Test Selection and Prioritization techniques is shown in Figure 1.

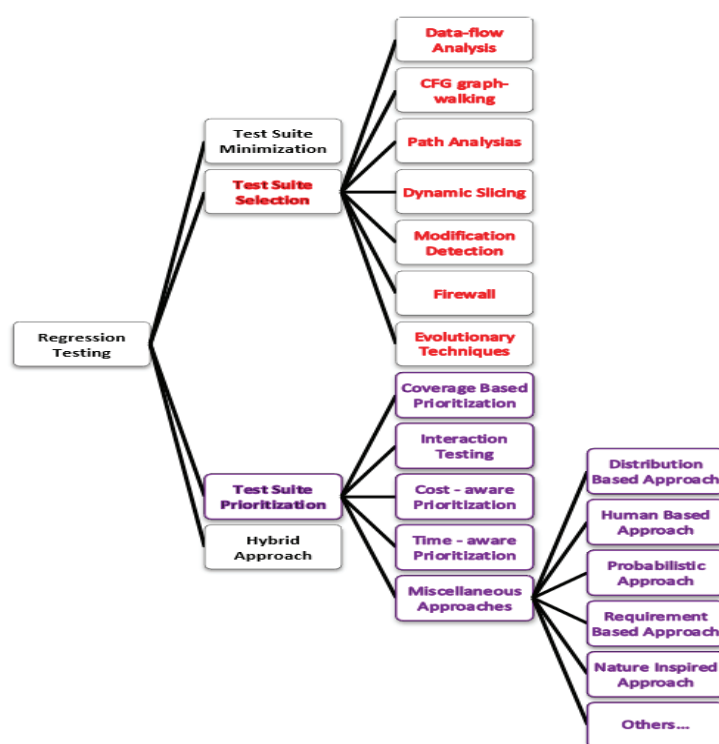


Figure 1. Taxonomy of Regression Test Selection and Prioritization.

Test Case Prioritization (TCP) [2] helps in ordering test cases based on certain criteria such as maximization of fault coverage in order to reduce the testing cost and speed up the delivery of the modified product. Regression test prioritization techniques try to reorder a regression test suite based on decreasing priority. The priority of test cases is established using some predefined testing criterion. The test cases having higher priority are preferred over lower priority ones in the process of regression testing [3]. Prioritization can be achieved on the basis of one or more objectives such as faults, code coverage, cost of execution, etc. Code coverage is a metric that measures the degree of coverage of the program code by a particular test suite [3]. The improved performance of the prioritized test suite in terms of reduced run time is a forever goal in this ever-time-constrained world. Furthermore, intelligent prioritization would mean that more and more faults are exposed as early as possible. Thus, the multi-objective goal is to deliver a regression test suite with the capacity to reveal maximum faults in the least possible time. The prioritization problem can thereby be stated as follows:

“Given: A test suite, T , the set of permutations of T , PT , and a function from PT to real numbers, f : $PT \rightarrow R$. To find $T' \in PT$ such that $(\forall T'')(T'' \in PT)(T' \neq T'')[f(T') \geq f(T'')]$.”

Test suite minimization or reduction seeks to minimize the number of test cases in it by removing the redundant ones. These concepts of reduction and minimization are interchangeable as all the reduction techniques may be utilized in producing a provisional subset of the test suite, whereas minimization techniques permanently remove the test cases. More formally, the test suite minimization problem is stated as follows [4]:

“To find a representative set of test cases from the original test suite that satisfies all the requirements considered ‘adequate’ for testing the program”.

Test Case Selection (TCS), resembles test suite minimization; as both of them intend to select few and effective test cases from the original test suite. The significant difference is whether the selection criteria focus on modifications in the SUT or not. Test suite minimization is generally built on the basis of metrics such as coverage measured from a sole version of the PUT [5]. On the other hand, in TCS, test cases are picked up based on the relevance of their execution to the variations between the original and the modified versions of the SUT. While TCS techniques also aim to reduce the number of test cases, the majority of these techniques focus on modified parts of the software under test. This means that the emphasis is to identify the modified sections of the software under test. This requires statistical analysis of the functional view of the software under test.

ACO_TCSP algorithm maps the real-world ants to digital ants while solving the regression test selection and prioritization problem. This approach has been proposed and validated and found to be very promising in earlier studies. Although, the major drawback, as found in other ACO applications, was ACO_TCSP falling into the local optima due to over-exploitation and less exploration of the solution space. The current study tries to solve this problem by enhancing the search space (increasing the number of ants) and making the exploitation more elite or intelligent. This work undertakes the following multi-objectives to be achieved:

- (1) Highest fault coverage achieved i.e., maximum faults should be caught by the test suite, and
- (2) The least possible test case execution time.

Hence, Enhanced ACO_TCSP has been proposed to achieve these two objectives, and the technique has been validated on four benchmark programs.

ACO is an established and well-tested optimization approach. There are a huge number of more recent optimization algorithms such as butterfly optimization algorithm (BOA), water optimization algorithms (WOA), Trees social relations optimization algorithm (TSR), red deer optimization algorithm (RDO), and many more. These have not been chosen for the current research as it tries to improve an already existing, well-established test case selection and prioritization algorithm that has already provided highly motivating results on prioritization.

The remaining paper is structured as follows: Section 2 presents the related work. Section 3 details the concepts of ACO and the limitations of the existing technique. Section 4 elucidates the proposed enhancements. Section 5 presents the implementation details of the enhanced technique by giving the algorithmic steps. Section 6 presents the experimental design. Section 7 presents the analysis of the results obtained. Sections 8 and 9 present the discussion and conclusion respectively.

2. Related Work

The way nature and its living beings co-exist in harmony is astonishing. The last two decades have witnessed extensive research in the area of developing and applying nature-inspired techniques for solving various combinatorial optimization problems [6]. Nature-inspired techniques such as ACO (Ant Colony Optimization) [7–9], GA (Genetic Algorithms), BCO (Bee Colony Optimization) [10,11], PSO (Particle Swarm Optimization) [12], NSGA -II (Nondominated Sorting Genetic Algorithm II) [13], Bat-inspired algorithm [14], flower pollination algorithm [15], cuckoo search algorithm [16,17], Cuscuta search [18], CSA [19], and hybrid approaches [20,21] (combining two or more different approaches) have already been applied to solve the regression test selection and prioritization problem.

There are other well-known regression testing techniques that deploy techniques such as fuzzy expert systems [22] and online feedback information [23].

All the ants have eaten your sweets? Powerful nature inspires us to build some artificial ants and achieve optimization goals. Inspired by the intelligent behavior of ants in food foraging, Dorigo gave a new soft computing approach and named it Ant Colony Optimization (ACO) [24]. This metaheuristic has already been applied to solve several non-polynomial time-solvable problems [25]. As given by Singh et al. [7], ACO can be applied to time-constrained test suite selection and prioritization. ACO variants have been applied rigorously by various researchers working in the area. Some of them include: epistasis based ACO [26], using ACO for prioritizing test cases with secure features [27] time-constraint-based ACO [28], history-based prioritization using ACO [29], prioritizing test cases based on test factors using ACO [30]. The widespread usage of ACO for test case selection and prioritization highlights the fact that ACO is a very prolific technique in the area. The improvements suggested in ACO [7,31,32] further ascertain that there are limitations associated with the technique. A survey of ACO in software testing [33] showed various limitations of ACO as identified by various researchers. This provided the key motivation for the authors for the present work. Thus, an improved version of the ACO technique for Test Case Selection and Prioritization has been proposed and implemented. The enhancement proposed has been validated in terms of: (1) no complexity overhead as compared to the previous approach, (2) improvement in correctness over the previous approach tested on four benchmark programs, and (3) APFD analysis against the traditional regression testing approaches [8]. The three enhancements proposed in the current manuscript have not been proposed for modification together in any of the above-mentioned related texts. Henceforth, the presented approach provides a new ACO approach having combined enhancements.

3. Ant Colony Optimization

This section gives a conceptual view of ACO and explains the limitations of the existing technique, thereby leading to the proposed enhancements in the subsequent section.

3.1. Concept

ACO is a popular metaheuristic method that provides a solution to many combinatorial optimization problems [34–36]. It is inspired by the concept of stigmergy, which is an indirect means of communicating with fellow members using the surrounding environment. This complex yet effective mode of communication is utilized by ants for food search. This working of ACO is shown in Figure 2. There are three possible paths (Path 1, Path 2, Path 3) from the ant nest to the food source. Ants cannot see, yet they magnificently coordinate within their colony for food search, with the use of a chemical substance known as pheromone. While foraging for food and fetching it back to their nests, ants keep laying the pheromone trail on the traversed path. The other ants thereby smell the maximum amount of pheromone and start following that path. The interesting fact to notice here is, that the ant on the minimum length path returns fastest, thus laying an additional amount of pheromone on their forward and return journeys. This enhances the probability of fellow ants taking this path. The ants taking this corresponding path will further drop more pheromones on that path. Therefore, attracting more fellow ants to take up this path. This process continues, and eventually, the whole colony of ants would converge to the minimum length path (here Path 1 in Figure 2).

The optimization algorithm based on the artificial behavior of ants works iteratively by the generation of an initial population. It then repeatedly endeavors to construct candidate solutions by exploration and exploitation of the search space. The exploration is guided by heuristic information and the experience gathered by the ants in the preceding iterations (well-known as ‘pheromone trails’) via a shared pool of memory. Using the concepts of using ACO, effective solutions have been achieved to many hard combinatorial problems [8]. The heuristic function (that marks the quality of the candidate solution

during the construction phase), and the pheromone values (signifying the information collected over different iterations) are the two major components of the working of ACO as a solution to combinatorial problems.

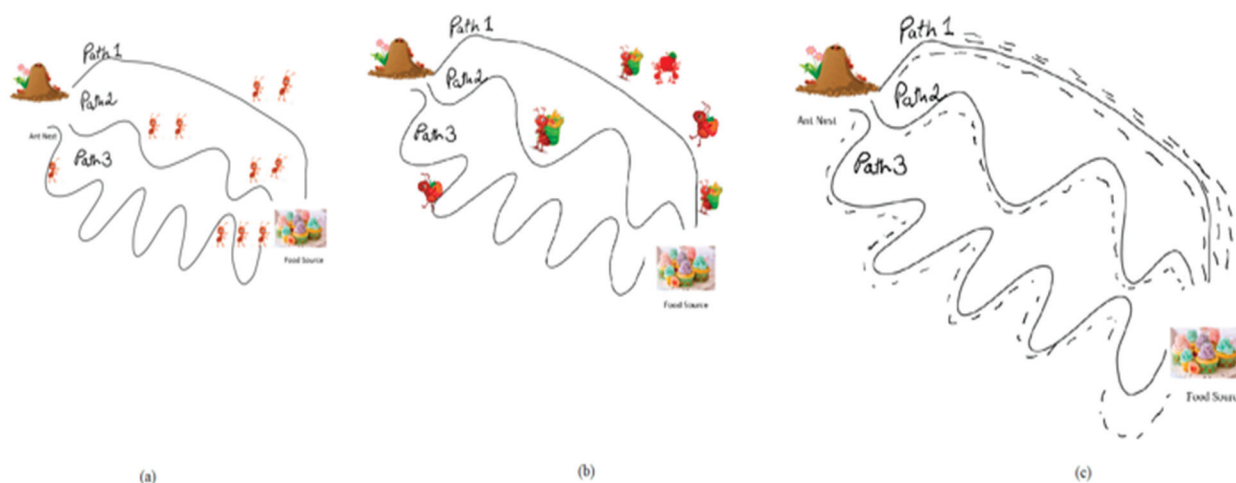


Figure 2. Pheromone on various paths. (a) Before ants begin searching, (b) while ants are searching for food, (c) finally, when all the ants follow the shortest path to the food source—the shortest path from the nest to the food source.

3.2. Limitations of ACO

ACO has been magnificently used to yield effective solutions to various optimization problems, yet its convergence has not been proven. This is also justified by the ACO approaches being approximation approaches. The general limitations include: slow convergence speed [33] and the problem of subsiding in the local optimal solution [33]. The reason behind it may be that the ACO updates the pheromone on the basis of the present promising path, and subsequently, after certain iterations, the amount of pheromone on this path rises substantially, while the pheromone for the probable worth path is frail. All the ants are inclined to this promising path and it becomes extremely tough to skip this path. Hence, the chances of marking the optimal solution being the locally optimal one, are increased.

Practically it is not viable to re-run the whole original test suite for building confidence in the modified software during regression testing. Thus, it becomes crucial to use selection and prioritization for the tests. Current research undertakes the highest fault coverage achieved in the least possible test case execution time as the objectives. ACO for test selection and prioritization was proposed in [7], implemented as ACO_TCSP in [8], and analyzed in [31,37]. From the analysis conducted in [31,37], some limitations in the ACO_TCSP technique were discovered:

- (1) Lower Termination condition (TC) values lead the ACO_TCSP to fall into the local optima problem by restricting exploration of newer paths.
- (2) Around 5% of the time, the ACO_TCSP could explore the optimum path, but due to the lack of pheromone that could be deposited, it is not the best path, i.e., the algorithm ended by meeting set TC criteria.
- (3) Lower TC value directly implies a lower number of iterations. Hence, changes are made in the algorithm to reach the set TC value as delayed as possible.

4. Proposed Enhancements

The major problem was falling into local optima due to a lack of explored search space and over-exploitation at an early stage. To solve this problem, it has been tried to expand the explored search space in each iteration while maintaining elite or intelligent exploitation to confirm as early convergence as possible to the optimal solution. Balancing increased

exploration with the intelligence introduced into the manner of exploitation should ensure that the run times are not adversely affected, while the solutions found are improved to avoid local optima problems. The proposed enhancements have been elaborated as follows:

4.1. Expand the Searched Space

In order to enhance the search space, we propose an increment in the number of exploring ants, which was earlier equal to the number of test cases in previous implementations. The enhanced algorithm asks the user to enter the enhancement factor from which the number of digital ants to be sent exploring paths per test case would be computed. Consequently, the overall number of ants grows multiple times the input entered by the user. If 'EF' is the Enhancement Factor input by the user, and '|TS|' represents the size of the original test suite, the total amount of ants (after enhancement) per program run can be considered to be:

$$\text{No. of Digital Ants} = (EF * |TS|) \quad (1)$$

The algorithm should now be able to overcome the problematic premature convergence of ants to a local optimal result. As now the ants will travel through additional paths in the initial iterations also, the number of paths to select the most promising one from for the current iteration correspondingly is multiplied by EF. Given as:

$$\text{Number of paths discovered in 1 iteration} = EF * |TS| \quad (2)$$

Enhancing the number of digital ants could directly affect the run time of our algorithm, but since the increase is by a constant enhancement factor (EF), it does not increase the complexity. Moreover, expanding the search space should lead to earlier convergence toward finding the optimal path. The same was also confirmed by the experimental results achieved (as shown in the next sections). So, the small constant time increase caused due to EF is counterbalanced by convergence at earlier iterations of the algorithm.

4.2. Elitism

In total, 5% of the earlier sample runs could find the optimum path if the overall best path computation neglects the maximum pheromone over the optimum best path found in any iteration. The calculation of the Global Best Path (GBP) was updated as:

$$GBP = \min\{\max \text{ pheromone path, Best_Path in any iteration}\} \quad (3)$$

This modification alone led to improved accuracy of ACO_TCSP by 4.1%. The elitist strategy for choosing the GBP over following the normal ACO approach is thus formulated. This step introduces eliteness in the process of exploiting already explored paths. This intelligent exploitation would lead to early convergence to the optimal solution. while at the same time, an increase in exploration would ensure the algorithm from falling into the local optima problems.

4.3. Modifying Total Time Calculation

The earlier algorithm took the MAX execution time (MET) of the local paths discovered in an iteration. The stopping criteria (TC) were compared with the MET added in every iteration. The modification to resolve this issue is recalculating MET using the execution time corresponding to the Local Best Path (LBP) discovered in every iteration.

$$MET = MET + \text{ExecTime of LBP} \quad (4)$$

The modification leads to an increased number of iterations for which the ACO_TCSP would run before reaching the set TC value. It leads to increased exploration and increased intelligent exploitation of paths and hence more chances of discovering the optimal path. It is of utmost importance to make sure that modifying MET or enhancing the search space does not adversely affect the execution time of the improved algorithm. Hence, the same

was taken care of by ensuring that the algorithm is time bounded by the user-entered TC value and cannot run after MET has reached TC. Hence, the Enhanced ACO_TCSP algorithm ensures better results without an increase in the run time of the process.

5. Implementing the Enhancements

Given ' TS ' (Test Suite) with $|TS|$ test cases in it, the selection and prioritization problem is specified as follows:

Find the subset ' S ' of test suite ' TS ', so as to have ' m ' test cases ($m < |TS|$, $S \subseteq TS$). The subset ' S ' is selected with the aim of maximum fault coverage and prioritized on the basis of the minimum execution time taken to run the entire subset.

5.1. Problem Representation and Execution Steps

A mapping for the selection/prioritization problem as an undirected graph $G(N, E)$ having N and E as the set of nodes and edges correspondingly was performed. Test cases were mapped as the nodes of the graph. Here, ' w_i ' represents the cost of ' i^{th} ' edge in the graph. This cost of edges maps to the trail of pheromone deposited on the edges $e_i \in E$. Pheromone deposition was correspondingly mapped to the multiple objectives of (1) fault coverage ' f_i ' on the current path achieved within (2) time constraint, ' Z_i '. Originally, the pheromone or cost on all the edges is nil.

- MAX represents the overall Time Constraint.
- Z stands for an intermediary variable used for TC calculations.
- $\{a1 \times EF, a2 \times EF, \dots, an \times EF\}$ represents a set of artificial ants, where 'EF' is the enhancement factor.
- S1 to S ($n \times EF$) be the sets in which ($n \times EF$) ants maintain the record of the covered test cases.

The execution steps of Enhanced ACO_TCSP are detailed below:

Step 1: Initialization—Generating $N \times EF$ ants to be sent for exploring the solution space. All other parameters are initialized accordingly.

Step 2: Exploration—Once generated, the ants begin searching in all random directions initially. As they move from one test case to another, their paths and killed faults are constantly updated on the way.

Step 3: Pheromone deposition—Once a probable solution has been found by each ant in the current iteration, pheromones are deposited on the most effective path (causing minimum time to run the test cases).

Step 4: Iterating—Steps 1 to 2 are repeated over and over till the stopping TC is achieved. Now there are many iterations and the most effective path from every iteration has been found.

Step 5: Finding GBP—Out of all the paths discovered during all the repetitions, GBP is the one with the least running time. This represents the selected and prioritized set of test cases as the final answer.

These steps are depicted in Figure 3.

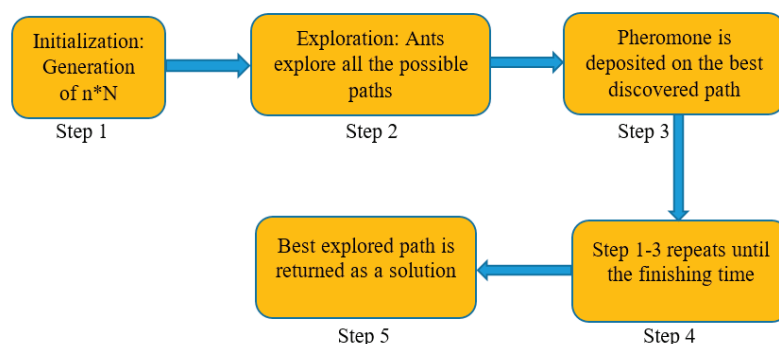


Figure 3. Working of Enhanced ACO_TCSP for Test Case Selection and Prioritization.

5.2. Modified Algorithm and its Complexity

Aiming to ascertain that the cost of applying the technique is no more than performing the complete regression testing, complexity has to be calculated. The complexities of all the sequential stages are added to compute the total complexity of the new proposed Enhanced ACO_TCSP (Algorithm 1) (Enhancements from [7,8]).

Algorithm 1: Enhanced ACO_TCSP

Stage-1

1.Initialization

	Statement-wise complexity
Set $W_i = 0$	 1
Set $TC = MAX$ (User Defined)	 1
.....	 1
Set $EF =$ Enhancement Factor (User defined)	 1
Set $Z_x = 0$ for all $x = 1$ to $(N \times EF)$	 $N \times EF$
Create $N \times EF$ artificial ants $\{a_1, a_2, \dots, a(N*EF)\}$	 $N \times$
EF	
... $S_1, S_2, \dots, S(N \times EF) = NULL$	 $N \times EF$

Stage-2

2. Do // LOOP 3—outermost

For $x = 1$ to $(N \times EF)$ // LOOP 2	 Runs $N \times EF$ times
$S_x = S_x + \{t_{xEF}\}$ // Initial test case for ant a_x is	 1
$tmpT = t_{xEF}$ // the starting vertex for ant a_x on the graph	 1

Do // Loop 1—Innermost loop

$tt = \text{Call_select_test_case}(a_x, tmpT)$	 6
$tUpdate\ S_x = S_x + \{t\}$	 1
$tZ_x = Z_x + t_{exec_time}$	 1
$tmpT = t$	 1

 While (total faults are covered) // Max N times (no of all Test Cases)

 // Loop 1 (Innermost loop)—ENDS

 EndFor // Loop 2 ENDS

$minTime = \min \{Z_x\}$	 $N * EF$
$currTime = CurrTime + minTime$	 1
$Z_x = 0$, for all $x = 0$ to $N \times EF$	 $N \times EF$
Update pheromone on all edges of bestPath with $minTime$	 $n - 1$
Evaporate $k\%$ pheromonee from each edge	 $n - 1$
$Last_Path = S_x$ for k of the last iteration of Loop2	 1
$S_x = null$ for all x	 EF

While ($TC \geq currTime$) // End of Loop 3

// MODULE select_test_case is presented below (As taken from [7,8])

select_test_case ($x, next_node$)

{

 If(max_pheromene edge out of all the edges from next_node to a node 'k'

 ($W[nxtnode][k]$) does not exist in $S[i][1 \dots N]$)

 Then

 Return 'k'

 Else

 Select a random edge $[next_node, k]$, from next_node to node 'k' not existing in $P[i][1 \dots N]$, from all edges having next max ($W[nxtnode][h]$)

 Return 'h'

 }

Stage 1 represents the initialization stage. The cost of the edges is initialized to be nil. EF is read from the user and corresponding digital ants are created. A set of paths and

other variables are also initialized as above. The total complexity of Stage 1, say S_1 , can be computed by adding the individual complexity of all the sequential statements from above as follows:

$$S_1 = 1 + 1 + 1 + (n \times EF) + (n \times EF) \quad (5)$$

$$S_1 = 3 + 3 \times (n \times EF) \quad (6)$$

The exploration of various paths keeping in mind the multiple objectives of fault coverage and budget time happens in this stage. Over the various repetitions of the do-while loop, pheromone updates occur, causing the next digital ants to exploit the already discovered potential paths. The complexity of Stage 2, say S_2 , can be calculated stepwise from the above individual statement complexities as shown below:

$$\text{Complexity of Loop 1} = n \times (6 + 1 + 1 + 1) = 10 \times n \quad (7)$$

$$\text{Complexity of Loop 2} = n \times EF \times (2 + (10 \times n)) = 2 \times EF \times n + 10 \times EF \times n^2 \quad (8)$$

$$\text{Complexity of Loop 3} = \text{Const} \times (3 \times EF + 2 \times (n - 1) + 2 + (2 \times EF \times n + 10 \times EF \times n^2)) \quad (9)$$

(Taking 'C' to be some positive constant)

$$S_2 = C \times (10 \times EF \times n^2 + 2 \times EF \times n + 2 \times n + 3 \times EF) \quad (10)$$

Stage-3

Stage 3 takes constant time steps to find the best out of the potential best paths according to the pre-decided multi-objectives. This stage takes (9—constant time instructions) as the worst-case complexity

$$S_3 = 9$$

Thus, the **Total Complexity** of the improved algorithm can be calculated as:

$$S = \text{Complexity of (Stage 1 + Stage 2 + Stage 3)} \quad (11)$$

$$S = S_1 + S_2 + S_3 \quad (12)$$

$$S = 3 + 3 \times (n \times EF) + C \times (10 \times EF \times n^2 + 2 \times EF \times n + 2 \times n + 3 \times EF) + 9 \quad (13)$$

$$S = C \times 10 \times EF \times n^2 + C \times 2 \times EF \times n + 3 \times EF \times n + C \times 2 \times n + C \times 3 \times EF + 12 \quad (14)$$

$$S \leq C \times 10 \times EF \times n^2 + C \times 2 \times EF \times n + 3 \times EF \times n + C \times 2 \times n + C \times 3 \times EF + 12 \times EF \quad (15)$$

(taking $C_1 = C \times 3 + 12$, to be some other positive constant)

$$S \leq C \times 10 \times EF \times n^2 + C \times 2 \times EF \times n + 3 \times EF \times n + C \times 2 \times n + C_1 \times EF \quad (16)$$

$$S \leq C \times 10 \times EF \times n^2 + C \times 2 \times EF \times n + 3 \times EF \times n + C \times 2 \times n \times EF + C_1 \times EF \times n \quad (17)$$

(taking $C_2 = 3 + C \times 2 + C_1 + C \times 2$)

$$S \leq C \times 10 \times EF \times n^2 + C_2 \times EF \times n \quad (18)$$

$$S \leq C \times 10 \times EF \times n^2 + C_2 \times EF \times n \times n$$

(taking $C_3 = C \times 10 + C_2$)

$$S \leq C_3 \times EF \times n^2 \quad (19)$$

(N_a is also a positive constant, thus $C_4 = C_3 \times EF$, be another positive constant)

$$S \leq C_4 \times n^2 \quad (20)$$

Therefore, the overall complexity of the proposed can be approximated as n^2 , or in terms of Big-O notation, it is $O(n^2)$ or $O(|TS|^2)$. This is the same as the complexity of the

old ACO_TCSP algorithm. Thus, the proposed approach is found to be as time efficient as the earlier approach (ACO_TCSP).

6. Experimental Design

The proposed approach (Enhanced ACO_TCSP) has been developed in C++. The organization of the code included ten modules and five global functions. The time for executing the code was computed with the help of a clock() function (available in the file “time.h”). The code was implemented on Macbook Pro 2011 model hardware, macOS 10.12 Sierra operating system and the C++ runtime environment used to run the code was Code::Blocks.

6.1. Benchmark Programs

We used 4 programs as benchmark programs. P1, P2, P3, and P4 were chosen so that the cases for minimum (P8), maximum (P2), and medium (P1, P3) correctness found by the previous algorithm could be compared with the proposed algorithm. Concise depiction of the programs, the lines of code in them, the mutations induced, and the test suite execution times are listed in Table 1.

Table 1. Depiction of Benchmark programs.

Program	Program Title	Lines of Code (kloc)	No. of Mutations	Test Suite Size (TS)	Execution Time (ms)
P1	Coll_Admison	0.281	5	9	1.0532
P2	Triangle_sides	0.037	6	19	3.82
P3	Basic_calculator	0.101	9	25	0.825
P4	Railways Booking	0.129	10	26	1.77

6.2. Design

Each benchmark program was executed on 4 varying values of the *EF* (Enhancement Factor), along with 5 varying *TC* values. ACO being an approximation technique, does not yield the same result for the same set of inputs every time. Henceforth, 10 runs for each *EF* value for each program and each *TC* value were recorded ($10 * 4 * 5 = 200$ runs for each program). Every execution run returns the path found (may be optimal or non-optimal). The time constraints chosen as the termination criteria were chosen to be 200, 300, 400, 500, and 600 with forty runs of each program. Reduction in total execution time needed, resultant optimality of solution, and the correctness of the technique were obtained from the output data. We also found the percentage improvement in correctness using the new algorithm.

6.3. ACO Parameter Settings for Enhanced ACO_TCSP

To maintain the levels of exploration (finding new paths) and exploitation (using existing pheromone knowledge) five ACO parameters need to be set in order to produce the near optimum results. Various parameter settings have been suggested and used over the years [8,38,39]. On the basis of experimentation and results from [8] and the provided enhancements, the parameter settings used for Enhanced ACO_TCSP are described in Table 2.

Alpha α , and Beta β are the control parameters to keep in check the exploration versus exploitation balance of ACO algorithm. The values chosen were kept the same as used by many researchers and in [7,8]. 10% Evaporation rate was used to make sure that digital ants map to real ants, and that the pheromone does not keep on depositing, it evaporates at the end of each iteration. Additionally, +1 was tuned to be the amount of pheromone to be deposited on the edges of the best path. Q_0 is a constant used for calculating the amount of pheromone to be evaporated joint with the evaporation rate. The value of N_a was updated

as per the proposed enhancements. To counter-balance the enhanced exploration, the value of τ has been updated to ensure elitism and q_0 ensures convergence. The rest of the parameter settings are the same as those used in the earlier version [8].

Table 2. Parameter Values used for Enhanced ACO_TCSP.

Na	N * EF
A	1.0 (Parameter to control exploration level)
B	0.2 (Parameter to keep exploitation in check)
P	0.1 (Evaporation Rate of pheromone—10% as used by majority researchers)
T	1.0 (Amount of Pheromone to be deposited on best path edges)
q0	1.0 (chosen constant used on calculation of pheromone to be evaporated)

7. Result Analysis

We now depict and explain the results of the implementation of the proposed enhanced technique here.

7.1. Execution Time of Paths Discovered

The details about the path and their corresponding execution time for four open-source programs have been given here. The combined results are shown in the next section.

The Enhanced ACO_TCSP was repeated for 10 Runs with EF values taken as 1, 2, 3, and 4 for program P1 using five values of TC . The averaged output of 10 runs for four values of the EF is summarized in Table 3. The details of the rest of the programs have been omitted due to redundancy and space constraints. However, the graphs depicting the results of the execution of all four programs have been depicted in the next section.

Possible best paths discovered in each of the 40 (5 TC , 4 EF) runs of Enhanced ACO_TCSP were compared for their execution times for all 4 test programs and are depicted graphically in Figure 4. Unlike previous results, a very slight decreasing trend is observed for increasing values of EF . It can be thus derived from here that an increased number of ants (proportional directly to the EF) leads to decreased execution time for the possible best paths. The exception for P2 is that the possible best path was found every time due to multiple possible best paths possible in the problem.

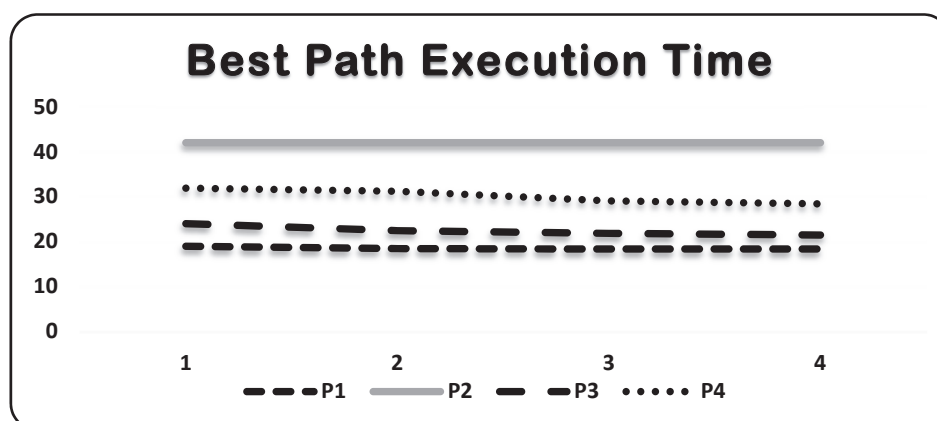


Figure 4. Best Path Execution Time of Enhanced ACO_TCSP versus EF .

Table 3. Enhanced ACO_TCSP on P1 with varying *TC* & *EF*.

RUN	Na	Running Time (sec)	Best Path	Total Exec. Time	Best P Exec. Time	No. of Test Cases Covered	No. of Iterations	Best P Found at Iteration	Optimal P Found at Iteration
TC = 200									
ants	1	0.054945	8	82.10127	18.796	2	10.7	2.4	2.25
ants	2	0.054945	9	82.3225	18.563	2	10.8	3.8	4
ants	3	0.049451	10	82.54895	18.33	2	10.9	3.2	3.2
ants	4	0.054945	10	82.54373	18.33	2	11	2.3	2.3
avg		0.053571		82.37911	18.50475	2	10.85	2.925	2.945946
TC = 300									
ants	1	0.043956	6	80.87737	20.096	2.1	15.3	3.2	3.666667
ants	2	0.06044	9	82.23752	18.647	2	16.2	4.5	4.555556
ants	3	0.054945	10	82.54373	18.33	2	16.9	1.8	1.8
ants	4	0.065934	10	82.53329	18.33	2	17	1.7	1.7
avg		0.056319		82.04798	18.85075	2.025	16.35	2.8	2.8
TC = 400									
ants	1	0.049451	8	81.96406	18.946	2	20.8	7.2	7.5
ants	2	0.054945	10	82.54373	18.33	2	22	3.4	3.4
ants	3	0.06044	10	82.53851	18.33	2	22	1.3	1.3
ants	4	0.071429	10	82.52808	18.33	2	22	1.8	1.8
avg		0.059066		82.39359	18.484	2	21.7	3.425	3.289474
TC = 500									
ants	1	0.054945	9	82.3225	18.563	2	26.5	6.2	6.555556
ants	2	0.054945	10	82.54373	18.33	2	27.4	3.7	3.7
ants	3	0.06044	10	82.53851	18.33	2	27.8	2.2	2.2
ants	4	0.071429	10	82.52808	18.33	2	27.4	1.3	1.3
avg		0.06044		82.4832	18.38825	2	27.275	3.35	3.358974
TC = 600									
ants	1	0.054945	10	82.54373	18.33	2	32.2	5.6	5.6
ants	2	0.065934	10	82.53329	18.33	2	32.5	5.7	5.7
ants	3	0.06044	10	82.53851	18.33	2	32.8	4.2	4.2
ants	4	0.076923	10	82.52286	18.33	2	33	1.8	1.8
avg		0.06456		82.5346	18.33	2	32.625	4.325	4.325

The average number of iterations over 40 runs each, observed in the case of *TC* and *EF* are presented in Figure 5. A steady surge in the number of iterations needed by Enhanced ACO_TCSP with an increase in input *EF* values can be observed. This is promising and encouraging because this implies that there is a reduction in the execution time of possible best paths earlier in the iterations. This causes more iterations to take place before reaching the stopping *TC*. Thereby, higher values of *EF*, indicate a rise in the number of iterations taken by Enhanced ACO_TCSP for the execution of the regression test suite under consideration, for the generation of the selected and prioritized test suite as resultant.

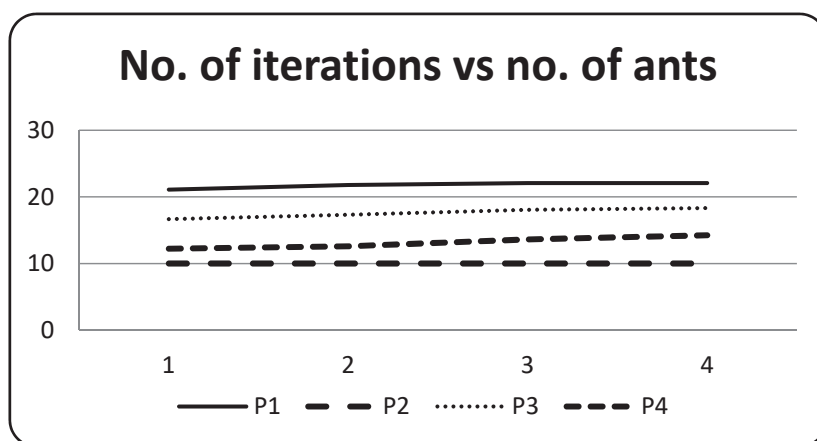


Figure 5. Average No. of Iterations v/s No. of ants sent.

The average iteration number at which the optimum path is found by the Enhanced ACO_TCSP execution out of 40 runs and its variation with different *EF* (no. of ants) is shown in Figure 6 (tabulated in the last column of Table 2 also). A steady fall in the iteration number capable of revealing the optimum path with the rising *EF* is clearly observable. P2 portrays an exception by revealing the optimum path in first iteration for all 200 test runs each. Consequently, by enhancing *EF*, the likelihood of locating the optimum path arises at even earlier stages of the improved technique. This fall in the ‘accurate iteration number’ (iteration that reveals the optimal path for the first time is called the ‘accurate iteration’), with the increase in *EF* is the couple effect caused by the improved exploration search space (that now avoids falling into local optima that could lead to inaccurate iterations), and elite exploitation of the earlier explored paths.

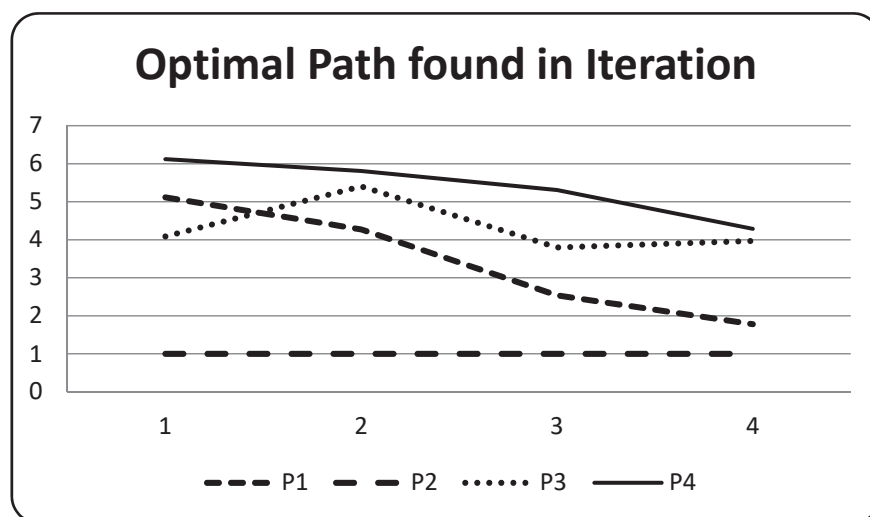


Figure 6. Averaged iteration number finding the Optimal path v/s EF.

The fact that ACO provides an optimum reduction of execution time has been witnessed in the previous section. The results have been averaged over 40 runs with four varying values of *EF* on the benchmark programs. The graphical analysis is presented in Figure 7. There is a slight rise in the reduction of execution time with the rise in *EF*. The observations are a result of the enhancements carried out in the technique. Better results have been yielded even at lesser values of *TC*, and are well-adjusted by the enhanced number of ants.

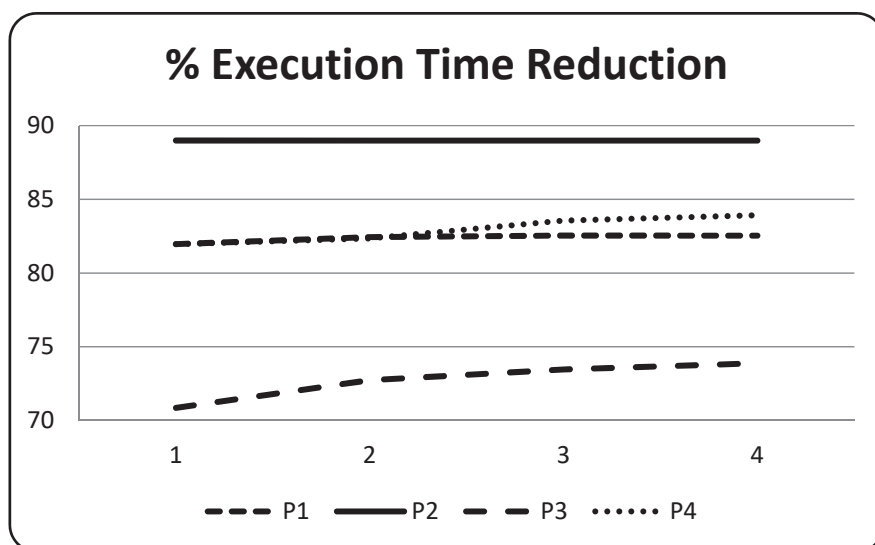


Figure 7. % execution time reduction for Enhanced ACO_TCSP selected test cases v/s EF.

7.2. Correctness of the Technique

This sub-section tries to prove the correctness of the technique, both theoretically as well as experimentally using the results of our empirical evaluation.

Theoretically, the improved algorithm ensures that a path is completed if all faults are found or all test cases have been visited. This ensures that the algorithm would stop within the computed complexity. In addition, the pheromone is then deposited on the best path from each iteration, ensuring exploitation of already found paths. ACO is a randomized approximation approach. Hence, the authors do not claim that the final test suite has minimum execution time. However, it is definitely found that the best APFD or fault coverage would be achieved by the final test suite. Experimentation will prove how many times ACO results in minimum execution time as well.

In order to prove the correctness of the proposed technique, Figure 8, Figure 9, and Table 4 shall be used to depict the correctness of the proposed work.

The percentage correctness of the Enhanced ACO_TCSP versus the EF has been picturized graphically in Figure 8. A very motivating and clear observation is the rise in the correctness achieved with the rise in EF (no. of ants). This validates the improvement and enhancements made in the prior ACO technique, which now is not falling to the local optima problems.

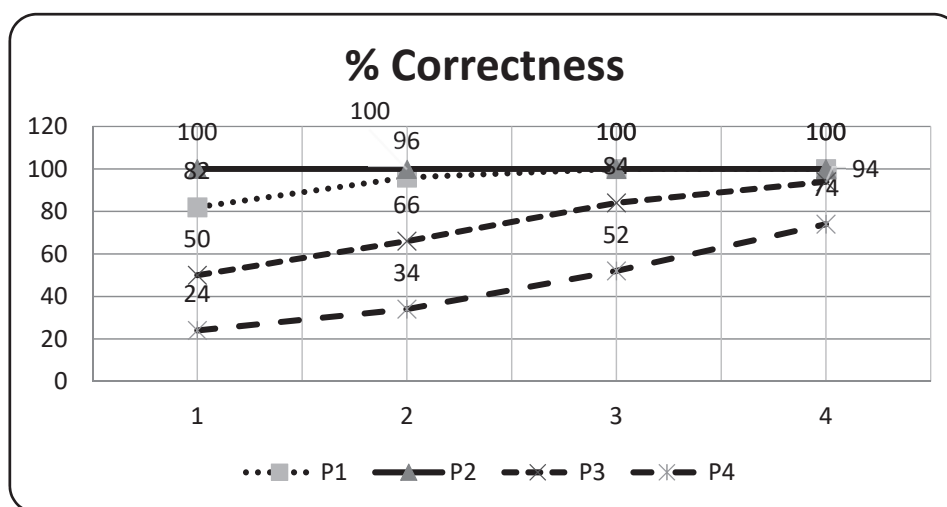


Figure 8. Percentage Correctness Achieved by Enhanced ACO v/s EF.

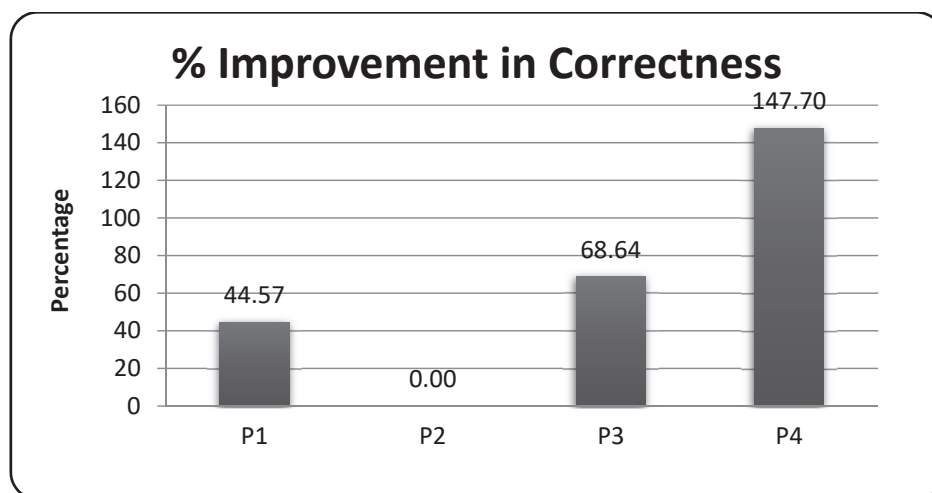


Figure 9. Percentage Improvement in Correctness Achieved by Enhanced ACO.

Table 4. Percentage improvement in Correctness Achieved by Enhanced ACO_TCSP.

Prog. No.	OLD % Correctness	NEW % Correctness	% Improvement
P1	65.714	95	44.565846
P3	100	100	0
P6	47.143	79.5	68.635853
P8	18.571	46	147.69802

Figure 9 depicts the improvement of percentage correctness achieved by the Enhanced ACO_TCSP approach. The achieved percentage improvement is nearly 50% in general, while for the earlier worst case of P4 using ACO_TCSP, an improvement of 147.7% is achieved using Enhanced ACO_TCSP. This clearly proves that the proposed and improved algorithm is better than the old ACO_TCSP algorithm.

7.3. APFD and Statistical Analysis

In order to obtain APFD (Average Percentage of Faults Detected), we calculated the area below the plotted line using a graph plot drawn between % of faults detected with the number of test cases needed. The notations used for calculation of APFD are:

‘TS’: the test suite containing the set of ‘|T|’ test cases,

‘F’: the set of ‘|F|’ faults revealed using ‘TS’.

For prioritization of test suite ‘TS’, let TF_i denote the priority order of the initial test case revealing the i th fault. The APFD for ‘TS’ can be obtained from the following equation:

$$APFD = 1 - \frac{TF_1 + \dots + TF_m}{|T| * |F|} + \frac{1}{2 * |T|} \quad (21)$$

Although, APFD is the popularly used criteria for the evaluation of the techniques used in the prioritization of test cases. Maximization of the APFD is not the objective of test case prioritization techniques. Maximization of APFD is a possibility when it is known in advance which faults are killed by a given test suite, thereby implying that the execution of entire test cases is already completed. Then, there would be absolutely no need for prioritization of test cases. APFD is thus needed after the task of prioritization for the evaluation of the prioritization technique.

The Enhanced ACO_TCSP orderings achieved for the four sample programs have been empirically evaluated (with respect to: No order, Random order, Reverse order, and optimum order of the test cases). These approaches are evaluated using APFD. Figures 9–12

depict the results obtained. It is evident that ACO attains results similar to that of optimum ordering, and has been shown to outweigh the reference techniques in terms of % of fault coverage achieved.

From Figure 10, it can be inferred that the same APFD value of 72.22% is achieved for the optimum and the ACO ordering for P1. Similar results are yielded for other programs also, as depicted in Figures 11–13.

The best APFD results achieved using the improved ACO algorithm, as depicted above, ensure that not only the maximum faults are covered in the entire test suite, but also maximum faults are revealed at earlier stages of running the prioritized test suite. These further motivate us to use the Enhanced ACO_TCSP algorithm.

In order to further ascertain the efficiency of our proposed work, we statistically analyzed the performance of existing ACO and the proposed Enhanced ACO_TCSP approach. We use the Independent Two Sample *t*-Test for statistical analysis. The objective here is to inspect if the proposed technique is more efficient than its earlier version. For this, we apply a *t*-test to compare % correctness and % time reduction achieved in the case of the four benchmark programs.

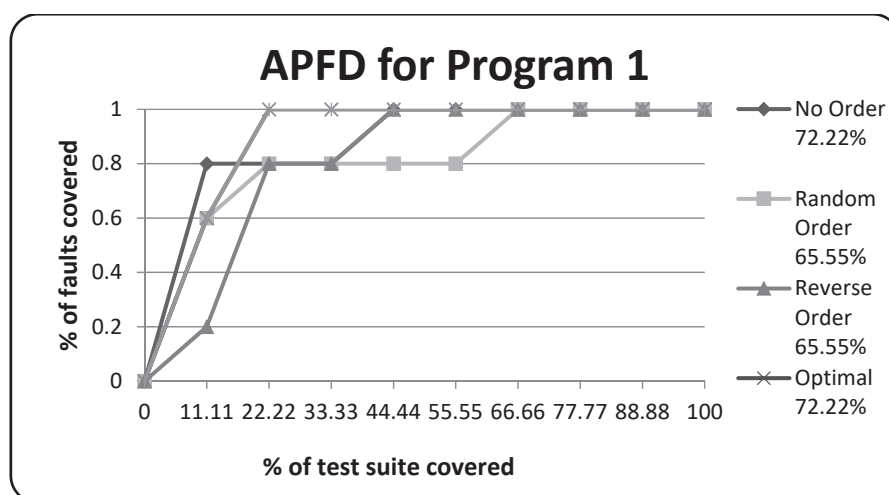


Figure 10. APFD for P1.

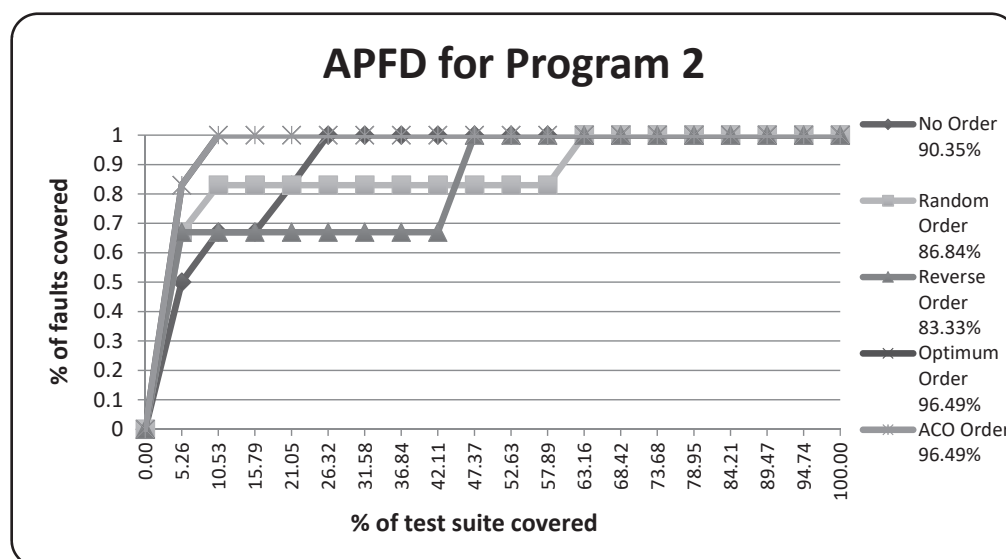


Figure 11. APFD for P2.

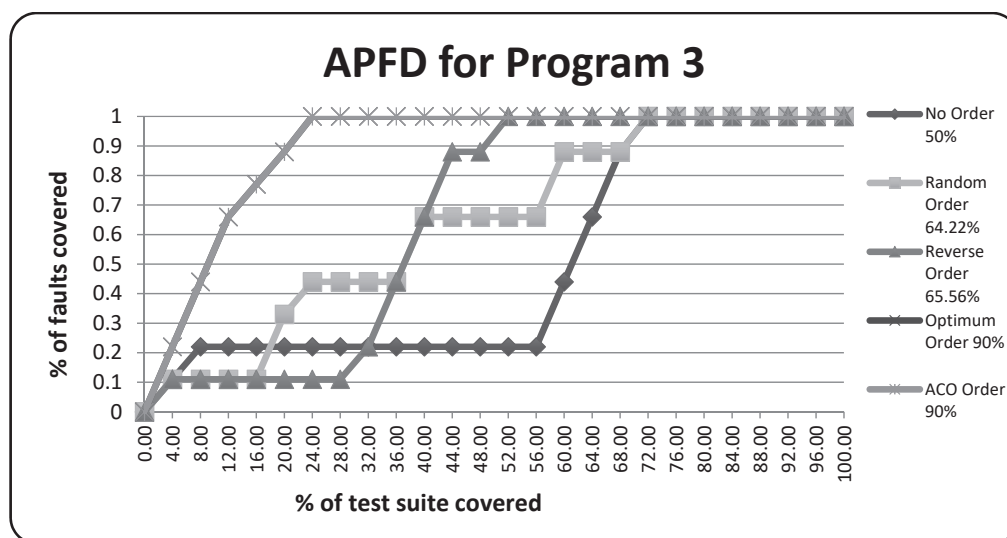


Figure 12. APFD for P3.

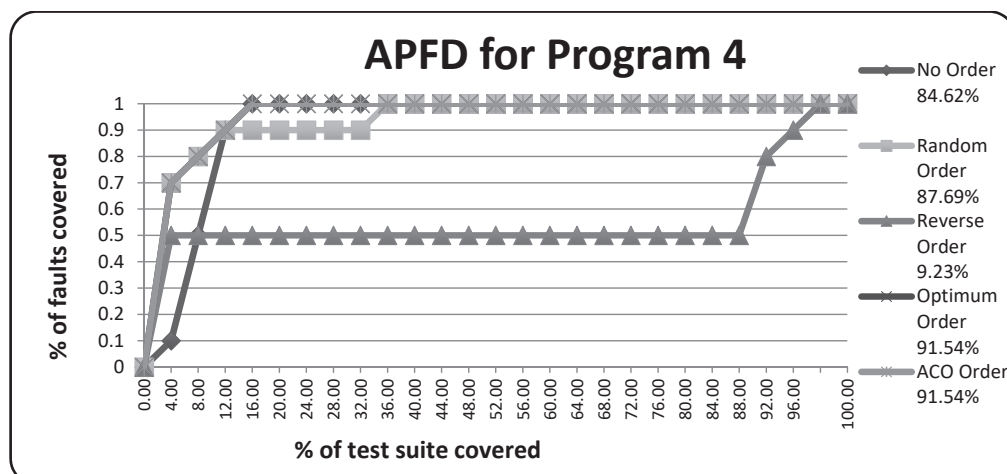


Figure 13. APFD for P4.

The Null Hypothesis (H_0) for the t -test is taken as:

(H_0 .) There is no significant improvement in the performance of Enhanced ACO_TCSP (in terms of % correctness and % time reduction) as compared to its earlier version (ACO_TCSP) for test data selection and prioritization.

To ascertain our claim, proving our research hypothesis and thereby rejecting the H_0 , the outcome of statistical analysis was conducted using Python language. The outcomes have been tabulated in Table 5 and depicted in Figures 14 and 15 below.

Table 5 can be observed to find that the probability of H_0 being true is significantly less (p -value < 0.5). Thereby we can reject the H_0 (Null Hypothesis) and state that our research hypothesis is true, that is, our proposed technique exhibits significantly improved performance over its earlier version.

Table 5. p -value obtained after applying an independent two-sample t -test.

p -value obtained for % correctness	0.05896016215814072
p -value obtained for % time reduction	0.11581073519437922

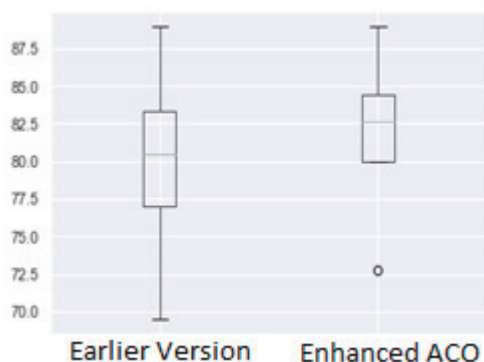


Figure 14. Comparison of % Correctness achieved.

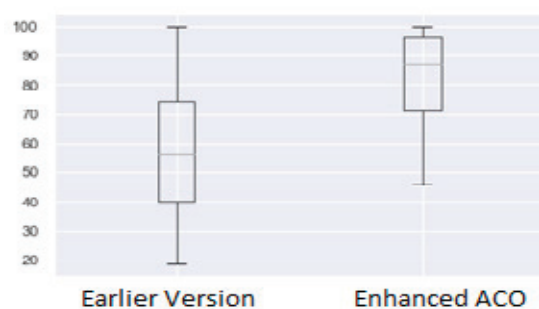


Figure 15. Comparison of % time reduction achieved.

The results of the *t*-test can be graphically analyzed using the following box plots:

Encouragingly, Figures 13 and 14 help us to unambiguously show the significant improvement achieved in the case of Enhanced ACO_TCSP in terms of % correctness and % time reduction for test case selection and prioritization. This provides a concrete validation to the proposed enhancements and motivates us to further use the enhanced approach in the field.

8. Discussion

In this paper, an enhancement of ACO for improving the test case selection and prioritization technique proposed by Singh et al. [7] has been developed and validated on four benchmark programs. Moreover, a comparison with five traditional prioritization techniques has been accomplished using APFD. The results achieved are encouraging owing to the following reasons:

- (1) The proposed technique results in the minimization of the test suite as the EF increases,
- (2) The running time is substantially reduced, and with the rise in EF, it is further reduced,
- (3) The precision of results achieved is encouraging for most of the test programs,
- (4) The percentage improvement in correctness is very high compared to the previous technique,
- (5) A comparison of the Enhanced ACO_TCSP prioritized test suite with No Order, Reverse Order, Random Order and Optimal Order prioritized test suites using APFD has been carried out. The results yielded APFD values for ACO that are equivalent to the optimum values (values that have maximum possible fault exposure in minimum possible time). The effect of the enhancement factor for different values of EF depicted motivating observations validates the Enhanced ACO_TCSP against the old approach [31]. The time reduction for the chosen resultant test suite by ACO was found to be almost the same for varying values of EF. This is due to the balancing provided by the increase in the number of ants for the new algorithm. The resultant test suite thus obtained potentially provides fast fault coverage.

The investigation of the usage of enhancement factor for different values of EF steers one to the following substantial leads:

1. A higher number of possible best paths are found at increased EF.
2. The selected best path is about the same for low and high values of EF.
3. The iteration number for convergence of Enhanced ACO_TCSP reduces with the increase in EF, this validates more exploration at the initial stages of the algorithm also.
4. Enhanced ACO_TCSP tends to yield optimum results at higher values of EF.

In addition to the above, statistical validation of the Enhanced ACO_TCSP has also been conducted. An Independent Two Sample *t*-Test has been performed to examine the correctness and execution time improvement achieved. Box Plots have also been used to represent the same. To affirm the validation of Enhanced ACO_TCSP, the *t*-test produced excellent results. All the aforementioned observations indicate that the proposed Enhanced ACO_TCSP technique yields promising solutions and exhibits better results than the existing technique in terms of solution correctness. The paper contributes to the literature by presenting the Enhanced ACO_TCSP approach and providing detailed enhancements and their validation on four benchmark programs without compromising the complexity of the algorithm. This can be easily re-implemented and fruitfully used by researchers for selecting and prioritizing test cases in the future.

9. Conclusions and Future Scope

The enhanced ACO_TCSP proposed in this work enhances the search space and makes the exploitation elite. The experimental results of the proposed approach on four benchmark programs were found to validate the enhancements. The average accuracy of the Enhanced ACO_TCSP was found to improve by over 30% over the original ACO_TCSP approach. Furthermore, the optimal paths converged at earlier iterations. All this could be achieved without an increase in execution time. This was ensured by the stopping criteria of TC (time-constraint) entered by the user. In addition to this, APFD analysis also proved the earlier exposure of faults achieved in comparison with the traditional prioritization approaches. Hence, as in the real world, increasing the size of the ant colony ensures more exploration, and elitism ensures intelligent exploitation of the already discovered paths. As in real ants, these ensure finding the optimal path with earlier convergence. Henceforth, this paper presents and validates the Enhanced ACO_TCSP approach for solving regression test selection and prioritization. As a part of future work, we can implement newer techniques [40] and empirically evaluate them for test case selection and prioritization. The proposed enhancements in ACO_TCSP proposed in this work can be applied and tested on many more metaheuristics that have been applied in the area of test case selection and prioritization for better efficiency in terms of results, without increasing the time complexity of these techniques. The limitations of the proposed work are as follows. The results have been experimentally evaluated on small codes. They can be experimentally evaluated in future work. The limitations associated with ACO as a methodology are also applicable to our proposed technique; however, the most prominent limitation of ACO getting stuck in local minimum has been averted using our technique due to the proposed enhancements. We have tried to imitate the behavior of real ants in our work; however, the intuitive behavior of the real ants cannot be incorporated even in the Enhanced ACO_TCSP technique.

Author Contributions: Conceptualization, S.S. and N.J.; methodology, S.S. and N.J.; software, S.S., N.J., K.S., G.D. and S.M.; validation, R.G. and B.S.; formal analysis, M.N., S.N.M. and N.R.P.; investigation, S.S. and N.J.; resources, S.S. and N.J.; data curation, S.S. and N.J.; writing—original draft preparation, S.S. and N.J.; writing—review and editing, S.S., N.J., K.S., G.D. and S.M.; visualization, S.S. and N.J.; supervision, B.S.; project administration, R.G. and B.S.; funding acquisition, M.N., S.N.M. and N.R.P. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Data Availability Statement: No data was needed in the work related to this research. The online sources from which the programs under test were downloaded re mentioned in the manuscript.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Bajaj, A.; Abraham, A.; Ratnoo, S.; Gabralla, L.A. Test Case Prioritization, Selection, and Reduction Using Improved Quantum-Behaved Particle Swarm Optimization. *Sensors* **2022**, *22*, 4374. [CrossRef] [PubMed]
2. Ansari, A.; Khan, A.; Khan, A.; Mukadam, K. Optimized regression test using test case prioritization. *Procedia Comput. Sci.* **2016**, *79*, 152–160. [CrossRef]
3. Kumar, S.; Ranjan, P. ACO based test case prioritization for fault detection in maintenance phase. *Intern. J. Appl. Eng. Res.* **2017**, *12*, 5578–5586.
4. Mohapatra, S.K.; Prasad, S. Finding Representative Test Case for Test Case Reduction in Regression Testing. *Int. J. Intell. Syst. Appl.* **2015**, *7*, 60. [CrossRef]
5. Noemmer, R.; Haas, R. An evaluation of test suite minimization techniques. In *Software Quality: Quality Intelligence in Software and Systems Engineering*; Springer: Vienna, Austria, 2019.
6. Sharma, B.; Hashmi, A.; Gupta, C.; Jain, A. Collaborative Recommender System based on Improved Firefly Algorithm. *Comput. Y Sist.* **2022**, *26*, 2. [CrossRef]
7. Singh, Y.; Kaur, A.; Suri, B. Test case prioritization using ant colony optimization. *ACM SIGSOFT Softw. Eng. Notes* **2010**, *35*, 1–7. [CrossRef]
8. Suri, B.; Singhal, S. Implementing Ant Colony Optimization for Test Case Selection and Prioritization. *Int. J. Comput. Sci. Eng.* **2011**, *3*, 1924–1932.
9. Wu, L.; Huang, X.; Cui, J.; Liu, C.; Xiao, W. Modified adaptive ant colony optimization algorithm and its application for solving path planning of mobile robot. *Expert Syst. App.* **2023**, *215*, 119410. [CrossRef]
10. Bajaj, A.; Sangwan, O.P. A systematic literature review of test case prioritization using genetic algorithms. *IEEE Access* **2019**, *7*, 126355–126375. [CrossRef]
11. Li, F.; Zhou, J.; Li, Y.; Hao, D.; Zhang, L. Aga: An accelerated greedy additional algorithm for test case prioritization. *IEEE Trans. Softw. Eng.* **2021**, *48*, 5102–5119. [CrossRef]
12. Jatana, N.; Suri, B. Particle swarm and genetic algorithm applied to mutation testing for test data generation: A comparative evaluation. *J. King Saud Univ.-Comput. Inf. Sci.* **2020**, *32*, 514–521. [CrossRef]
13. Chaudhary, N.; Sangwan, O. Multi Objective Test Suite Reduction for GUI Based Software Using NSGA-II. *Int. J. Inf. Technol. Comput. Sci.* **2016**, *8*, 59–65. [CrossRef]
14. Öztürk, M.M. A bat-inspired algorithm for prioritizing test cases. *Vietnam. J. Comput. Sci.* **2018**, *5*, 45–57. [CrossRef]
15. Dhareula, P.; Ganpati, A. Flower Pollination Algorithm for Test Case Prioritization in Regression Testing. In *ICT Analysis and Applications: Proceedings of ICT4SD 2019*; Springer: Singapore, 2020; Volume 93, pp. 155–167. [CrossRef]
16. Srivastava, P.; Reddy, D.P.K.; Reddy, M.S.; Ramaraju, C.V.; Nath, I.C.M. Test Case Prioritization Using Cuckoo Search. In *Advanced Automated Software Testing: Frameworks for Refined Practice*; IGI Global: Hershey, PA, USA, 2020; pp. 113–128.
17. Kaushik, A.; Verma, S.; Singh, H.J.; Chhabra, G. Software cost optimization integrating fuzzy system and COA-Cuckoo optimization algorithm. *Int. J. Syst. Assur. Eng. Manag.* **2017**, *8*, 1461–1471. [CrossRef]
18. Mann, M.; Sangwan, O.P. Test case prioritization using Cuscuta search. *Netw. Biol.* **2014**, *4*, 179–192.
19. Jatana, N.; Suri, B. An Improved Crow Search Algorithm for Test Data Generation Using Search-Based Mutation Testing. *Neural Process. Lett.* **2020**, *52*, 767–784. [CrossRef]
20. Panwar, D.; Tomar, P.; Singh, V. Hybridization of Cuckoo-ACO algorithm for test case prioritization. *J. Stat. Manag. Syst.* **2018**, *21*, 539–546. [CrossRef]
21. Yoo, S.; Harman, M. Using hybrid algorithm for Pareto efficient multi-objective test suite minimisation. *J. Syst. Softw.* **2010**, *83*, 689–701. [CrossRef]
22. Xu, Z.; Gao, K.; Khoshgoftaar, T.M.; Seliya, N. System regression test planning with a fuzzy expert system. *Inf. Sci.* **2014**, *259*, 532–543. [CrossRef]
23. Zhou, Z.Q.; Sinaga, A.; Susilo, W.; Zhao, L.; Cai, K.-Y. A cost-effective software testing strategy employing online feedback information. *Inf. Sci.* **2018**, *422*, 318–335. [CrossRef]
24. Dorigo, M.; Maniezzo, V.; Colormi, A. The Ant System: An Autocatalytic Optimizing Process. In *Technical Report TR91-016*; Politecnico di Milano: Milano, Italy, 1991.
25. Chaudhary, R.; Agrawal, A.P. Regression Test Case Selection for Multi-Objective Optimization Using Metaheuristics. *Int. J. Inf. Technol. Comput. Sci.* **2015**, *7*, 50–56.
26. Bian, Y.; Li, Z.; Zhao, R.; Gong, D. Epistasis based aco for regression test case prioritization. *IEEE Trans. Emerg. Top. Comput. Intell.* **2017**, *1*, 213–223. [CrossRef]

27. Vescan, A.; Pinte, C.M.; Pop, P.C. Solving the test case prioritization problem with secure features using ant colony system. In Proceedings of the International Joint Conference: 12th International Conference on Computational Intelligence in Security for Information Systems (CISIS 2019) and 10th International Conference on European Transnational Education (ICEUTE 2019), Seville, Spain, 13–15 May 2019; Springer: Cham, Switzerland, 2019.
28. Suri, B.; Singhal, S. Understanding the effect of time-constraint bounded novel technique for regression test selection and prioritization. *Int. J. Syst. Assur. Eng. Manag.* **2015**, *6*, 71–77. [CrossRef]
29. Noguchi, T.; Washizaki, H.; Fukazawa, Y.; Sato, A.; Ota, K. History-Based Test Case Prioritization for Black Box Testing Using Ant Colony Optimization. In Proceedings of the IEEE 8th International Conference on Software Testing, Verification and Validation (ICST), Graz, Austria, 13–17 April 2015. [CrossRef]
30. Ahmad, S.F.; Singh, D.K.; Suman, P. Prioritization for Regression Testing Using Ant Colony Optimization Based on Test Factors. In *Intelligent Communication, Control and Devices*; Springer: Berlin, Germany, 2018; pp. 1353–1360. [CrossRef]
31. Singhal, S.; Jatana, N.; Suri, B.; Misra, S.; Fernandez-Sanz, L. Systematic literature review on test case selection and prioritization: A tertiary study. *Appl. Sci.* **2021**, *11*, 12121. [CrossRef]
32. Suri, B.; Singhal, S. Evolved regression test suite selection using BCO and GA and empirical comparison with ACO. *CSI Trans. ICT* **2016**, *3*, 143–154. [CrossRef]
33. Suri, B.; Singhal, S. Literature survey of Ant Colony Optimization in software testing. In Proceedings of the CONSEG, CSI Sixth International Conference On Software Engineering, Indore, India, 5–7 September 2012. [CrossRef]
34. Kavitha, R.; Jothi, D.K.; Saravanan, K.; Swain, M.P.; Gonz  les, J.L.A.; Bhardwaj, R.J.; Adomako, E. Ant colony optimization-enabled CNN deep learning technique for accurate detection of cervical cancer. *BioMed Res. Int.* **2023**, *2023*, 1742891. [CrossRef]
35. Singhal, S.; Suri, B. Multi objective test case selection and prioritization using African buffalo optimization. *J. Inf. Optim. Sci.* **2020**, *41*, 1705–1713. [CrossRef]
36. Singhal, S.; Jatana, N.; Dhand, G.; Malik, S.; Sheoran, K. Empirical Evaluation of Tetrad Optimization Methods for Test Case Selection and Prioritization. *Indian J. Sci. Technol.* **2023**, *16*, 1038–1044. [CrossRef]
37. Suri, B.; Singhal, S. Analyzing test case selection & prioritization using ACO. *ACM SIGSOFT Softw. Eng. Notes* **2011**, *36*, 1–5.
38. Dorigo, M.; Di Caro, G.; Gambardella, L.M. Ant Algorithms for Discrete Optimization. *Artif. Life* **1999**, *5*, 137–172. [CrossRef]
39. Chen, X.; Gu, Q.; Zhang, X.; Chen, D. Building Prioritized Pairwise Interaction Test Suites with Ant Colony Optimization. In Proceedings of the 2009 Ninth International Conference on Quality Software, Jeju, Korea, 24–25 August 2009; pp. 347–352. [CrossRef]
40. Zhu, A.; Xu, C.; Li, Z.; Wu, J.; Liu, Z. Hybridizing grey wolf optimization with differential evolution for global optimization and test scheduling for 3D stacked SoC. *J. Syst. Eng. Electron.* **2015**, *26*, 317–328. [CrossRef]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

Article

Proposal of a Framework for Evaluating the Importance of Production and Maintenance Integration Supported by the Use of Ordinal Linguistic Fuzzy Modeling

Ronald Díaz Cazañas ¹, Daynier Rolando Delgado Sobrino ^{2,*}, Estrella María De La Paz Martínez ¹, Jana Petru ³ and Carlos Daniel Díaz Tejeda ¹

¹ Faculty of Mechanical and Industrial Engineering, Central University “Marta Abreu” de Las Villas, 5^{1/2} Km Camajuaní Road, Santa Clara 54830, Cuba; ronaldcdc@uclv.edu.cu (R.D.C.); estrella@uclv.edu.cu (E.M.D.L.P.M.); cadtejeda@uclv.edu.cu (C.D.D.T.)

² Faculty of Materials Science and Technology in Trnava, Institute of Production Technologies, Slovak University of Technology in Bratislava, 917 24 Trnava, Slovakia

³ Faculty of Mechanical Engineering, VSB—Technical University of Ostrava, 70800 Ostrava, Czech Republic; jana.petru@vsb.cz

* Correspondence: daynier_sobrino@stuba.sk

Abstract: Over the years, the integration of Production Management and Maintenance Management has gained significant attention from the scientific community due to its benefits for the company. When searching through the states of the art and practice, it is possible to understand that one of the main challenges for the integration is the lack of systematic, methodological, and scientific approaches and evaluation systems that lead companies into a successful implementation and a clear understanding of the benefits and drawbacks of the process. This paper introduces an original framework that conducts the processes of evaluation, weighting, and aggregation of set of novel indicators proposed by the authors. The main output of the proposal is an integral index that allows us to qualify, in a linguistic domain, the importance of the Production and Maintenance Management integration. At the same time, the proposed framework includes a methodology to evaluate the consensus of the experts, based on the use of linguistic terms with a membership function of the triangular type, which attempts to overcome some deficiencies of previous models identified by the authors in a detailed and complex analysis of the scientific literature. The proposed framework is applied in a plant of the Cuban mechanical industry. The results of this application are clearly presented and discussed, allowing us to verify and validate the proposal while also contributing to its ease of understanding and ultimately to the successful integration of the production and maintenance tasks in the given company.

Keywords: production management; maintenance management; production and maintenance management integration; ordinal linguistic fuzzy model; consensus evaluation

MSC: 90-10; 90-11; 03B52; 90C70; 90C90

1. Introduction

The success of any company requires an adequate level of integration between its systems and/or subsystems and processes. A particular case is the integration between the Production Management (PM) and Maintenance Management (MM) processes, given the potential benefits achieved in terms of cost reduction, increased availability, and performance when both subsystems fit and are managed with the required synergy [1–6].

In this sense the authors in [7] stated that some manufacturing objectives, such as inventory reduction, improved equipment reliability, increased productivity and quality, have led maintenance managers to operate proactively, coordinating their activities with the production function.

Similarly, authors like the ones in [8–10] pointed out that the MM and manufacturing strategy integration is vital to the same MM performance, which allows ensuring of the availability of the equipment, product quality, deliveries on time, and competitive prices. The work presented in [11] also states that the concept of maintenance must be reviewed periodically to take into account changes in the system and its environment; while the authors in [12] recognize that the perspective of the maintenance function will change depending on its interaction with the company, and thus the researchers and practitioners often have to deal with maintenance concept variations from one organization to another, and with the lack of a global solution easily adjustable to each case study.

At the same time, it is fair to say that the need to achieve an adequate level of integration between PM and MM may be more or less significant, depending on factors related to the design characteristics of the production subsystem and the effectiveness of the maintenance subsystem. In this regard, authors like the ones in [13,14] highlight that, in companies with a high component of labor and poorly integrated processes, the inverse nature of the Production and Maintenance functions (the operators use the equipment for producing and it has to be running while maintainers use it for restoring its condition while this usually has to be switched off) does not necessarily constitute a key issue given that the bottlenecks are not necessarily found on the machines. On the other hand, in complex industrial or service systems, in which there is significant automation and integration of processes such as, for example, in the automotive industry in Europe, the situation of the reverse activities of production and maintenance operators (operators require that equipment is upon on state, available for producing, meanwhile maintenance staff generally use it upon off mode for its restoring) is an increasingly important issue to take into account. In consequence, the need to adopt proactive maintenance policies and computer-assisted Maintenance Management systems is greater in large companies, with continuous production processes and high costs due to stoppages, but usually not in the case of small companies with process-oriented production flows. Similarly, the need to achieve integrated maintenance is more significant in those companies that wish to maintain control of their processes and flexibility at competitive levels.

Concerning this topic, the author in [15] points out that there are relationships between production technology and maintenance practices whose strength varies according to the complexity of each production environment. Similarly, the researchers in [11,16] also emphasize that the problem lies in the fact that companies with different competitive priorities just follow different maintenance strategies. The authors of this paper have no option but to identify themselves with both of the previous statements given that they see undeniable truth in both of them.

Despite all of this evidence on the importance, need, and benefits of PM—MM integration, there are only a few relevant studies in the states of the art and practice covering the topic from a perspective similar to the one pursued with this research. Even when none of these works has presented a clear method specifically designed to assess the importance of PM—MM integration within the company, it is still highly valuable, convenient, and necessary to examine their findings and conclusions in order to support a goal of this paper that lies in the proposal of a set of indicators leading to an index capable of evaluating the importance of PM—MM integration in a given company.

Based on this knowledge and partial conclusions, it is particularly important to mention an empiric research presented by [14], where the author focused on demonstrating that, when implementing a PM—MM integration strategy, companies had higher possibilities of obtaining a better overall performance. In addition, in this study, it was also concluded that the effects of maintenance approaches on the operational performance differed according to the given context or production instance. This author also analyzed the integration of maintenance within manufacturing from the perspective of two variables, i.e., hard integration and soft integration. The first of these referred to the technological support and computer-aided Maintenance Management systems that facilitate integration, while the second referred to the integration from the perspective of the relationships established

in the organizational structure, including human resources. His study covered a sample of 253 companies representative of the manufacturing sector in Sweden, and included indicators of the main maintenance dimensions, production context factors, and operational performance. Also, from a cluster analysis, three groups of companies were identified based on the type of maintenance strategy adopted. Those groups of companies differed in terms of size, process technology, and orientation of the manufacturing strategy.

A similar research approach was adopted by [15] to infer information about the relationship and correlation between production technology and the use of maintenance practices. The study revealed a significant and positive correlation between the decentralization of the maintenance workforce, with the technical complexity and interdependence between the stages of the production process. At the same time, the technical complexity also manifested a significant and positive correlation with the training and use of a professional maintenance staff, and a significant and negative correlation with the involvement of operators in maintenance tasks. No significant correlation was found between production technology and preventive maintenance practices, nor between the latter and technical variety.

In the case of [7], the authors developed a model to study the influence of the integration between maintenance policy and aggregate production planning over the total cost reduction. The study focused on three important elements, all of them related to cost impacts. Firstly, the integration of the maintenance policy with the aggregate planning of the production led to a reduction in the total cost. In addition, the effectiveness of the integration between production planning and maintenance planning proved to be more significant in productive environments where the financial consequences of failures were greater. Finally, the study also demonstrated that, when the maintenance cost increased in relation to the production cost, the savings associated with integration decreased as well.

Similarly, in [17], the authors developed a structural equation model to evaluate the relationship between the Total Productive Maintenance (TPM) philosophy and manufacturing performance (MP). As for the TPM, this was measured through variables related to autonomous maintenance and planned maintenance practices, while in the case of MP, it was measured using variables such as cost, quality, deliveries, and flexibility. After the implementation of the model, it was concluded that TPM had a significant and positive correlation with the variables that make up the MP construct.

In the case of [18], a framework to characterize the complexity presented in the MM within a productive environment was developed. The authors identified and made use of the following factors: availability of data and automated Maintenance Management systems, complexity and variety of manufacturing technologies, level of automation and integration of the process, Production Management system, level of knowledge of operators about the process, technology and maintenance, and the experience of maintenance personnel. Subsequently, in order to obtain an index able to reflect the complexity of MM in the company, a Likert scale was used to weight the factors according to the characteristics of the production system in the analyzed company. Although not as complete and complex as the proposals presented in this paper, the authors found the work in [18] to be useful and to work as a nice inspiration for our own proposals.

Similarly, the research in [16] presented an empirical study that evidenced that the way of managing maintenance was different between companies that emphasized different competitive manufacturing priorities. This was attributed to the fact that different business objectives required different levels of emphasis on maintenance related terms.

In [8], the authors developed a multi-attribute decision-making model for selecting the most appropriate maintenance system for the organization. The selection was based on the impact of each maintenance system on the performance of the organization. It was described from a set of criteria: the level of use of the equipment and the workforce, the reduction of defects of the product, of the process, and rework, the increase of average time between failures, and the reduction of the maintenance costs, the delays in the deliveries, the number of accidents, among others.

On the other hand, in [11], a framework that integrates the fundamental elements of the maintenance flexibility concept was proposed. According to these authors, business objectives decide on the maintenance strategy, and thus maintenance tasks must change in consonance with business objectives changes. These changes in maintenance tasks will in consequence define modifications in the same maintenance system and processes, which lead these authors to conclude a review of the infrastructure that supports the maintenance service, the information systems, and the education and training plans, is also required.

All of the elements mentioned above demonstrate the need to manage PM and MM processes in an integrated manner, and furthermore, they highlight the fact that the emphasis on the need for integration varies according to context. In this sense, the development of a method capable of evaluating how significant it is to maintain an adequate level of PM—MM integration, in a given company, would allow us to objectively demonstrate to the parties involved (mainly the production and maintenance managers) the need to deploy a decision-making process that is coherent and synergistic. All of it seeking to obtain the global optimum in the performance of the Production—Maintenance system, as a whole, instead of trying to generate a local optimum within each subsystem, which could be detrimental to the overall effectiveness and performance of the company.

Based on these considerations and detailed analysis of the relevant and related states of the art and practice, the present paper introduces a novel framework to direct the activities of evaluation, weighting, and aggregation of a group of indicators (attributes) as the base for obtaining an index that allows qualifying, in a linguistic domain, of the importance of PM—MM integration. At the same time, this framework includes a new methodology to evaluate the consensus of experts and uses linguistic terms with a triangular membership function. It is meant to overcome some deficiencies of previous models in which the index that was to characterize the consensus and did not adequately reflect the levels of proximity or distance between the linguistic terms emitted by an expert during the evaluation of attributes [19–21].

On the other hand, since the method selected for obtaining the weights together with the recursive character of the linguistic term aggregation operator used (LOWA operator) causes that the weight of the element to be added in each iteration i tends to be dominated by the one obtained in iteration $i - 1$, even if this element were to have a high real relative importance, the authors have decided to overcome this deficiency by introducing a parameter that quantifies the distance between linguistic terms into the expression determining the consensual relationship. Such a parameter has a triangular membership function. Furthermore, the authors also propose a new way of obtaining the components of the vector of weights used in the aggregation, so that in each iteration i the weight of the attribute to be aggregated reflects the real comparative importance of it with respect to the average relative importance of the attributes already aggregated in previous iterations.

The reminder of this paper is structured as follows: Section 2 presents a set of novel indicators to assess the different dimensions of the importance of PM—MM integration, while Section 3 details a novel framework designed to evaluate, weight, and aggregate the defined indicators. Given the qualitative nature of some of the proposed indicators, the evaluation process of these is based on the use of Ordinal Fuzzy Linguistic Modeling (OFLM), and after the aggregation phase, it is possible to obtain a general index that characterizes the importance of the PM—MM integration in the company under question, i.e., a Production–Maintenance Integration Index (PMII). Subsequently, Section 4 presents the results of the application of the proposals in a workshop of the Cuban mechanical industry, while at the end, Section 5 is dedicated to the conclusions and some suggestions for the development of future work on this subject.

2. Proposal of a Set of Indicators for Evaluating the Importance of the PM—MM Integration

Although in the states of the art and practice the authors could not find a proposal of indicators for specifically evaluating the integration level between PM and MM, the theoretical elements analyzed in the previous section as to the relationship Production—Maintenance serve as a strong base in the definition of such a set of indicators, by applying a deductive process. It is a proved fact that, in highly automatized plants, the importance of the PM—MM integration is more critical than in those intensive in manpower [11,13–16]. This reason leads us to think that the automatization level is an element that should be taken into account in the set of indicators. On the other hand, considering the fact that some authors point out that companies with different competitive priorities generally follow different maintenance strategies [7,11,15], then a competitiveness-driven approach is another of the indicators to be proposed.

A study presented in [7] showed that the integration between production and maintenance plans was more significant in production settings in which the failures' consequences were higher. This study also revealed that, as the maintenance cost increases in relation to the production cost, the savings derived from integration are reduced. If taking this fact into consideration, it is then useful to consider a group of elements that, at the floor shop level, characterizes and explains the failure consequences; likewise, it is also important to consider other elements linked with maintainability, which condition and determine the throughput, availability, and costs. In this sense, in terms of elements that highlight the need for the PM—MM integration, the authors of this paper also decided to consider others such as the impact of maintenance-related downtimes or stops over production goals; the rate of maintenance interventions that should be executed over a stopped machine; the complexity of reparations; and the level or extent of technical service required by the equipment.

On the other hand, as factors that decrease the need of the PM—MM integration, because of the reduction they lead to in terms of the failures' negative operational consequences, the authors of this paper have identified others, such as the existence of redundant equipment or alternative production lines [14], and the availability of resources for executing the maintenance activities [11].

From the organizational point of view, the implementation of a process-based management strategy requires an adequate coordination and integration between all the processes and systems of the company, being a particular case of those related to production and maintenance. For this reason, another of the proposed indicators is the recognition of the need to adopt process-based management.

According to the opinion of the authors of this paper and also of those in [13], the existence of common material resources for production and maintenance promises a higher PM—MM integration to eliminate the redundancy of information and informatic modules that manage these resources. For this reason, the proportion of material resources that are common to Production and Maintenance is another of the proposed indicators in this paper.

Based on all previous considerations, Table 1 presents a general proposal of unique and novel indicators that are the base and lead to the subsequent proposal of an index for evaluating the importance of the PM—MM integration in a company. Table 1 also presents a brief description of each of the indicators and the theoretical conceptions behind these.

The evaluation of each indicator will be either qualitative or quantitative and will depend on the own nature of these, given that indicators such as a competitiveness-driven approach could be hardly quantified while others such as the proportion of material resources that are common to Production and Maintenance can be easily expressed in numbers. In this sense, the research in [19] developed a model for the management of heterogeneous information in which data could be linguistic or numeric within the interval (0; 1). However, this work did not detail a way of normalizing quantitative indicators that could be expressed in a different interval or whose goal or ideal value could correspond to

the minimum or intermediate value within the interval. This gap is something that will be also addressed within the proposals in this paper.

Table 1. Proposed indicators for evaluating the importance of the PM—MM integration.

Indicators	Theoretical Conceptions
1. Automatization level	In processes highly depend on equipment, the need of integration is higher than in those manpower-based processes [11–16].
2. Competitiveness-driven approach	Companies with different competitive priorities follow different maintenance strategies [11,16].
3. Rate of interventions with stopped machine	If this results in a high value, the maintenance strategy should focus more on reducing the mean time between interventions. It requires higher coordination between both departments to take advantage of opportunity windows in the production program.
4. Complexity of the reparations	As maintainability decreases, the integration between both systems should increase to reduce the negative impact over availability.
5. Impact of maintenance stops over production goals	This has to be considered by the maintenance department in order to develop prevention or consequences' reduction strategies, definition of critical equipment, among others. According to [15], functional areas with low to no buffers should work in a coordinated way, otherwise small equipment interruptions can affect the whole system.
6. Proportion of material resources common to production and maintenance	This should be considered in order to eliminate problems caused by information redundancy. The higher its level, the higher the needs of integration in areas such as inventory and suppliers' management, design, and implementation of IT solutions in both systems, etc. [13].
7. Level of technical Service required by the equipment	If it is high, it demands higher integration, which can lead to development of an autonomous maintenance and/or outsourcing maintenance strategy.
8. Recognition of the need to adopt process-based management	As with the competitiveness-driven approach, the processes-based management approach forces the integration in the management of these processes, being a key and a particular case of the processes of PM and MM.
9. Resource availability for maintenance	It is one of the elements that supports maintenance quality [11]. If it results in a high value, the need for integration decreases. Otherwise, it indicates the need for a coordinated work in the search of strategies to reach the required availability.
10. Existence of redundant equipment	This decreases the tensions in reaching coordination between production and maintenance planning, because it avoids process interruptions by unexpected breakdowns or planned maintenance interventions that can overlap with production orders [14].

3. Framework for the Evaluation, Weighting, and Aggregation of Indicators Characterizing the Importance of the PM—MM Integration

This framework includes evaluation, weighting, and attribute aggregation modules, which allow us to obtain a general index that characterizes the importance of PM—MM integration in a company (PMII), see Figure 1. For the sake of a better understanding, each of the framework's steps will be described next:

3.1. Setting of the Goals

The goal values of the parameters that characterize the efficiency of the evaluating and weighting processes are set in this step.

3.2. Module I: Evaluation of the Attributes

This module uses the defined attributes (indicators) and the necessary information for their evaluation as input data. It includes the evaluation process and the quality analysis of this process itself, which is performed through the determination of a consensus degree or index. The output of this step is a set of evaluated attributes (indicators) in accordance with their current state and values in the company under study.

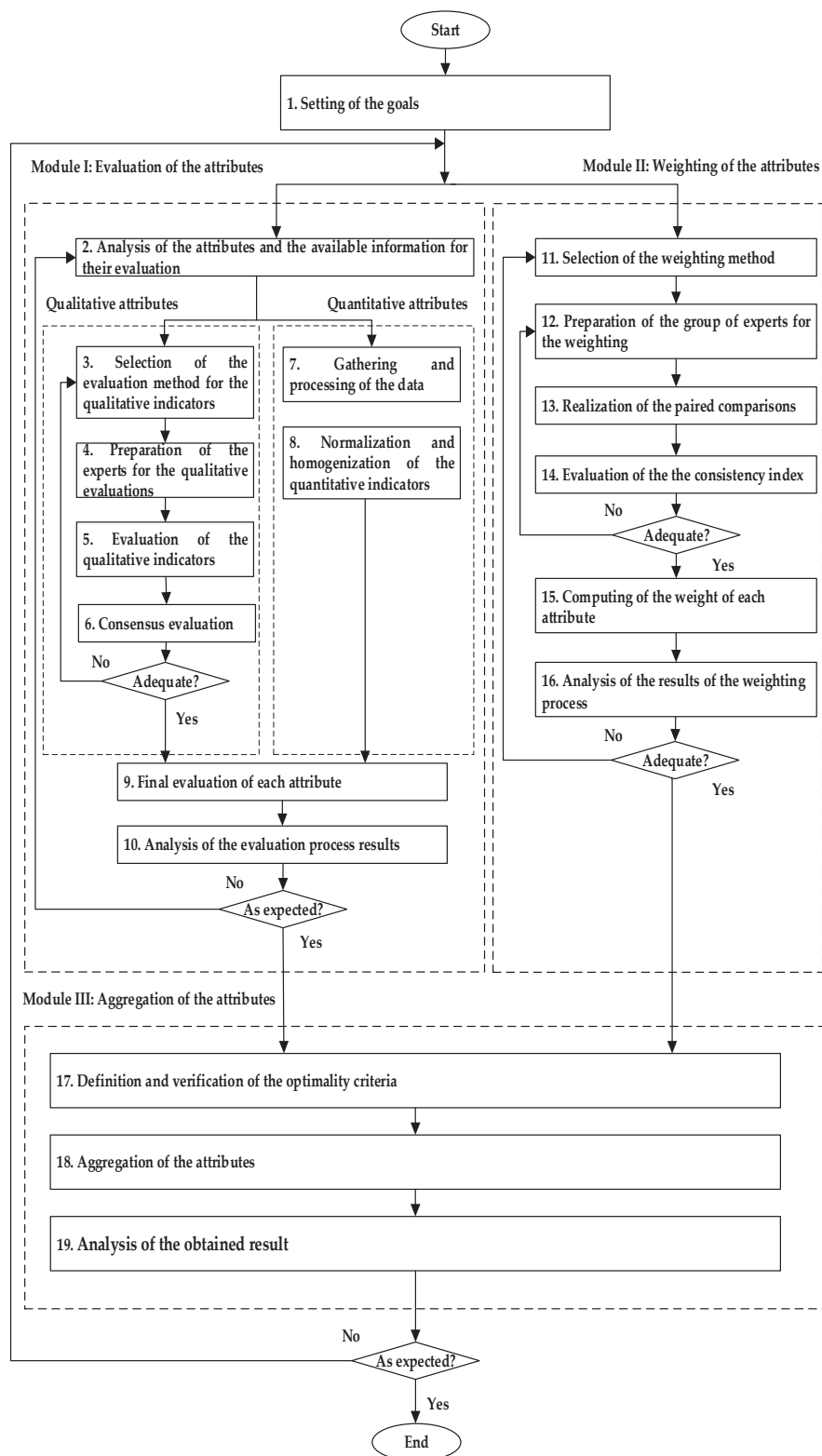


Figure 1. Framework for the evaluation, weighting, and aggregation of indicators characterizing the importance of the PM—MM integration.

3.2.1. Analysis of the Attributes and the Available Information for Their Evaluation

By considering and assessing the uncertainty level of the available information, it is possible to specify which indicators are to be evaluated in a qualitative way (linguistically), through the expert's judgment, and those which will be evaluated in a quantitative way.

3.2.2. Selection of the Evaluation Method for the Qualitative Indicators

The Ordinal Fuzzy Linguistic Modeling approach will be used as the evaluation method for the qualitative indicators. This decision is based on its capabilities for the treatment of the uncertainty and vagueness existing in the linguistic information provided by the experts [19–25]. The evaluation scale will be formed by a set of seven terms or linguistic labels, these are: null (n), very low (vl), low (l), medium (m), high (h), very high (vh), perfect (p), as shown in Figure 2. Readers of this paper are also kindly invited to check the content of [26], where these authors made a really complete and complex analysis on the Consensus Reaching Process for the modeling of experts' linguistic preferences.

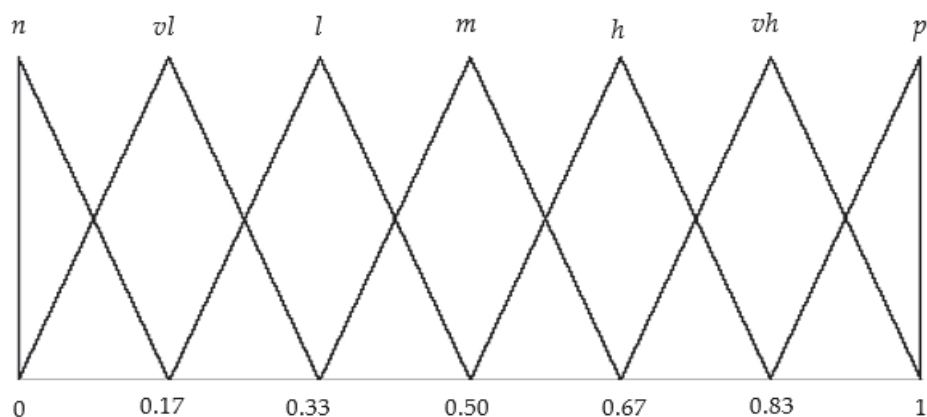


Figure 2. Adopted scale for evaluating the qualitative indicators.

For the definition of the linguistic terms' semantics, a triangular membership function of the type ($\mu_A(x)$) will be used, as it appears in Equation (1), see [24]. All of these terms will be represented through triangular fuzzy numbers with the parameters a_1 , a_2 , and a_3 , as is also shown in Figure 2.

$$\mu_A(x) = \begin{cases} 0 & \text{if } x < a_1 \\ \frac{x - a_1}{a_2 - a_1} & \text{if } a_1 \leq x \leq a_2 \\ \frac{a_3 - x}{a_3 - a_2} & \text{if } a_2 \leq x \leq a_3 \\ 0 & \text{if } x > a_3 \end{cases} \quad (1)$$

By setting S as the set of terms or linguistic labels we have the following:

S_0 = null (n), S_1 = very low (vl), S_2 = low (l), S_3 = medium (m), S_4 = high (h), S_5 = very high (vh), S_6 = perfect (p).

3.2.3. Preparation of the Experts for the Qualitative Evaluations

The content of each indicator will be explained to the experts. This is performed in order to warranty that each expert is able to establish an adequate match between the linguistic terms in the scale and the possible states of the indicator under assessment.

3.2.4. Evaluation of the Qualitative Indicators

Experts will assign to each indicator the linguistic term (label) of the scale S that, in accordance with their assessment, better and more objectively represents their current state in the company.

3.2.5. Consensus Evaluation

The authors in [20] proposed a good and well-accepted methodology for determining the consensus in a linguistic context, which is based on the calculation of one parameter called consensus relation (CR) and the subsequent application of a linguistic quantifier

over the obtained numerical value. However, although their approach and methodology allow us to evaluate the distance between the different assessments given by each expert and a collective opinion measure, it does not consider how close to each other the linguistic terms selected and assigned by the own expert can be, and for this reason, the CR would be the same in case of existing, within the group of experts, assessments of the type vl and vh , or h and vh , with these last two terms consecutive ones in scale where, logically, the consensus level would be expected to be higher.

In addition, when analyzing this research, it could be also noted that, when consecutive linguistic terms with cardinalities higher than 1 coexist, the consensus relation produces a low consensus level if the linguistic quantifier named “much” is used. This manifests a lack of coherency with the evaluation’s results if the linguistic terms’ proximity is taken into account, because of the fact of being consecutive values in the proposed scale.

Taking into consideration the above-mentioned facts, this section introduces the proposal of a methodology that overcomes the above-mentioned drawbacks or limitations. It includes the following steps:

1. To obtain a vector v_{si}^j for each indicator (attribute) j , whose components are the linguistic terms $s_i \in S$ established by the group of experts as a measure of the evaluation of the indicator j .
2. To obtain the quantity (G) of experts’ subgroups that is formed withing the group, in accordance with the coincidences in their evaluation of the attribute j :

$$G = \# \left(v_{si}^j \right) \quad \forall s_i \in S \quad (2)$$

where $\#$ represents the cardinality of the vector v_{si}^j .

3. For each linguistic term $s_i \in S$ of the vector v_{si}^j , to obtain the number of experts ms_{ij} that coincided in the assignation of it as the attribute j ’s evaluation measure.
4. To obtain the experts proportion (p_{si}) that coincided in assigning the linguistic label (term) s_i as the attribute j ’s evaluation measure:

$$p_{si} = \frac{ms_{ij}}{m} \quad (3)$$

where m refers the number of experts.

5. To calculate the consensus relation (CR) using the following expression:

$$CR_i = \begin{cases} 1 & \text{if } G = 1 \\ \sum_{j=1}^G p_{si}^2 & \text{if } G = 2 \text{ and the 2 linguistic terms are consecutive ones} \\ \left(\sum_{j=1}^G p_{si}^2 \right) \left(\frac{d_{min}}{\bar{d}} \right)^{\min(p_{sij})} & \text{if non - consecutive linguistic terms exist} \end{cases} \quad (4)$$

where:

CR_i : consensus relation reached by the experts in evaluating the indicator j .

p_{si} : expert proportion that coincided in assigning the term s_i as measure of the indicator j .

G : subgroup of experts formed in accordance with the coincidence in the utilized term for evaluating the indicator j .

d_{min} : normalized minimum distance between G consecutive linguistic terms of the scale.

$\bar{d}$: normalized mean distance between the non-consecutive terms (s_i) of the scale employed the experts for evaluating the indicator j .

Given that a triangular fuzzy number A , with parameters a_1 , a_2 , and a_3 , can be represented in accordance with the concept of the confidence interval of level α in the following way [27]:

$$A_\alpha = [a_1^\alpha, a_2^\alpha] = [a_1 + (a_2 - a_1)\alpha, a_3 - (a_3 - a_2)\alpha] \quad (5)$$

then, the normalized distance between two triangular fuzzy numbers A and B will be determined by Equation (6), as used in [27]:

$$d(A, B) = \frac{1}{2(\beta_2 - \beta_1)} \int_0^1 (|a_1^\alpha - b_1^\alpha| + |a_2^\alpha - b_2^\alpha|) d\alpha \quad (6)$$

where:

$d(A, B)$: normalized distance between the fuzzy numbers A and B both with a triangular membership function.

β_1 : $\min[a_{1A}; a_{1B}]$.

β_2 : $\max[a_{3A}; a_{3B}]$.

a^α and b^α : representation of the triangular fuzzy numbers A and B , respectively, based on the confidence interval concept of level α (α - cut).

$|a_1^\alpha - b_1^\alpha|$: distance to the left between the triangular fuzzy numbers A and B .

$|a_2^\alpha - b_2^\alpha|$: distance to the right between the triangular fuzzy numbers A and B .

For obtaining d_{min} , the distance among all the subsets of consecutive terms of size G that are generated starting of a subset of size seven is to be determined first. This has to be performed in accordance with the used evaluation scale (see Figure 2). Then, the one with the shorter distance is to be selected.

$$d_{min} = \min(d_{gi}) \quad \forall i = 1, \dots, N \quad (7)$$

where N refers to the number of subgroups of consecutive terms of size G that can be generated starting of seven terms, and d_{gi} the normalized mean distance existing among the G linguistic terms of the group i .

6. To select a linguistic quantifier that will represent the fuzzy majority concept.

A linguistic quantifier is a fuzzy subset Q that, for any value $r \in \mathbb{R}^+$ $Q(r)$ indicates the grade, level, or extent to which the value r satisfies the concept represented by the same Q . For more information on this, readers are kindly invited to check the vast study presented in [20].

7. Determination of the consensus level (CL_j) reached over each attribute j .

The equation to use is the following:

$$CL_j = Q^2(CR_j) \quad (8)$$

$Q^2(CR_j)$ can be obtained using Equation (9):

$$Q^2(CR_j) = \begin{cases} l_0 & \text{if } CR_j < a \\ l_i & \text{if } a \leq CR_j \leq b \\ l_u & \text{if } CR_j > b \end{cases} \quad (9)$$

With l_0 and l_u the minimum and maximum linguistic terms, respectively, of the employed linguistic terms set S , and a and b the minimum and maximum values of the quantifier's domain.

$$l_i = \sup l_q \in \{l_q\} \quad (10)$$

With

$$M = \left\{ lq \in L : \mu lq(CR_j) = \sup_{t \in I} \left\{ \mu_{lt} \left(\frac{CR_j - a}{b - a} \right) \right\} \right\} \quad (11)$$

j is the S linguistic terms set's cardinality.

8. Determination of the global consensus level (GCL) obtained in the evaluation of the attributes' set.

This GCL will be calculated by using the following equation:

$$GCL = Q^2 \left(\frac{\sum_{j=1}^n CR_j}{n} \right) \quad (12)$$

With n the number of attributes (indicators).

Next, the steps for evaluating the quantitative indicators are also presented.

3.2.6. Gathering and Processing of the Data

In this step, quantitative data are collected for evaluating this kind of indicator, in addition, a reference threshold value is also established for each indicator that allows evaluation of its significance.

3.2.7. Normalization and Homogenization of the Quantitative Indicators

The normalization process will take place using the next equation, which is also an original proposal within this paper:

$$Vn_j = 1 - \frac{|V_{r_j} - V_{m_j}|}{V_{r_j}} \quad (13)$$

Vn_j : indicator (attribute) j 's normalized value.

V_{mj} : current value of the indicator j .

V_{r_j} : reference value of the indicator j .

Subsequently, the normalized values should be homogenized so that they are expressed in the same scale, which appears in the previous Figure 2. For these purposes, the previously presented linguistic quantifier Q^2 will be applied over the values Vn_j .

3.2.8. Final Evaluation of Each Attribute

The final evaluation of each qualitative attribute will be obtained by aggregating the linguistic terms emitted by the experts as a measure of its evaluation, and this only if an adequate consensus level has been demonstrated. For such aggregation purposes, there are various aggregating operators that directly operate over the set of linguistic labels, see [21,23,24]. Various of these operators are based on the well-known Linguistic Ordered Weighted Averaging operator (LOWA) and others, such as LAMA. However, they all imply complex calculations that might restrict their application's efficiency. In this sense, for the purposes of the present research and based on the authors' long-term experience in this area, the sole application of the LOWA operator with some modifications and improvements to it, that are proposed later in this paper and help in obtaining the weight vector, seems to be sufficient. The mathematical apparatus and explanation behind this operator are shown as follows [23]:

$$\phi(a_1, \dots, a_m) = W \cdot B^T = C^m \{w_k, b_k, k = 1, \dots, m\} = w_1 \otimes b_1 \oplus (1 - w_1) \otimes C^{m-1} \{\beta_h, b_h, h = 2, \dots, m\} \quad (14)$$

where W is a weight vector $[w_1, \dots, w_m]$ such that $w_i \in [0; 1]$ and $\sum w_i = 1$;

$\beta_h = \frac{w_h}{\sum_{k=2}^m w_k}$, $h = 2, \dots, m$. $B = b_1, \dots, b_m$ is a vector associated to A such that $B = \sigma(A) = (a(\sigma_1), \dots, a(\sigma_m))$, being $a(\sigma_j) < a(\sigma_i) \forall i \leq j$.

Similarly, σ is a permutation defined over the set of terms A , while C^m is the m label convex combining operator, defined by the authors in [28].

For $m = 2$ we have the following equations:

$$C^2\{w_i, b_i, i = 1; 2\} = w_1 \otimes s_j \oplus (1 - w_1) \otimes s_i = s_k; s_i, s_j \in S(j \geq i) \quad (15)$$

k can be obtained using Equation (16):

$$k = \min\{T, i + \text{round}(w_1 \cdot (j - i))\} \quad (16)$$

With S the linguistic term set to be aggregated with cardinality $T + 1$.

Usually, the weight vector W is calculated by some quantifier, however, in this work, obtaining of the weights is proposed to be performed based on the cardinality of each of the terms to be aggregated, and this in accordance with the aggregation concept that it is commonly addressed and used for the majority operators without a quantifier [22]. This way we have:

$$\begin{aligned} w_i &= f_i(b_i, K, b_n) \\ &= \frac{\gamma_i^{\delta_{\min}}}{\theta_{\delta_{\max}} \cdot \theta_{\delta_{\max}-1} \cdot \theta_{\delta_{\min}+1} \cdot \theta_{\delta_{\min}}} + \frac{\gamma_i^{\delta_{\min}+1}}{\theta_{\delta_{\max}} \cdot \theta_{\delta_{\max}-1} \cdot \theta_{\delta_{\min}+1}} + \dots \\ &\quad + \frac{\gamma_i^{\delta_{\max}}}{\theta_{\delta_{\max}}} \end{aligned} \quad (17)$$

where:

b_i : i -th element of the term set to be aggregated, ordered in an increasing way in accordance with its cardinalities.

δ_i : Cardinality of the element i .

$$\gamma_i^K = \begin{cases} 1 & \text{if } \delta_i \geq K \\ 0 & \text{otherwise} \end{cases} \quad (18)$$

θ_i is calculated using Equation (19):

$$\theta_i = \begin{cases} (\text{number of elements with cardinality} \geq i) + 1 & \text{if } i \neq \delta_{\min} \\ \text{number of elements with cardinality} \geq i & \text{otherwise} \end{cases} \quad (19)$$

Because the aggregation process based on the LOWA operator is an iterative one, in which in each iteration element of the B linguistic term set is aggregated with the result obtained in the previous iteration, it happens that, in each iteration i , the term's weight to be aggregated tends to be dominated by the linguistic term's weight that was obtained as a result of the aggregation in the iteration $i-1$. This fact occurs because the term obtained as result in the iteration $i-1$ is to be aggregated in the iteration i with a weight value that constitutes the cumulated weight of all terms that have been aggregated until this iteration $i-1$.

In this regard, the authors of this paper propose a modification to the way of obtaining the weight vector β_h , where it is achieved that the weight value of each element reflects the relative importance previously calculated for it with respect to the mean relative importance of the elements that preceded it in previous iterations. The following expression (20) introduces the calculation method for this vector. This also constitutes a scientific added value in the form of an own proposal of the author's in this paper.

$$\beta_h = \begin{cases} \frac{w_h(m-h)}{w_h(m-h) + \sum_{i=h+1}^m w_i} & \text{if } h = 1, 2, \dots, m-1 \\ 1 - \beta_{m-1} & \text{if } h = m \end{cases} \quad (20)$$

where:

$\beta_h = [\beta_1, \beta_2, \dots, \beta_m]$ is a weight ordered vector associated to A such that: $w_{\beta_i} \geq w_{\beta_j} \forall i < j$.
 m : number of linguistic terms to be aggregated.

Now, considering the previous steps, the LOWA operator will be used, taking into consideration the previous weights (ϕ_F); while the reordering of the vector B will be now performed in a non-increasing way, in accordance with the new weights values (β_h) of the terms to be aggregated.

In the case of the quantitative attributes, the evaluation will be obtained directly from the previous step once the quantifier Q^2 is applied over the homogenized value (Vh_j).

3.2.9. Analysis of the Evaluation Process Results

In this step, the obtained evaluations will be presented to the experts with the aim of verifying their agreement with the results.

3.3. Module II: Weighting of the Attributes

This module uses as input information the operational characteristics of some of the used methods for the weights' calculation, the same nature of the attributes to be weighted, while it is also supported by the addition and use of some informatic application to ease the computing. The main processes developed in this module are the same weighting process and the evaluation of the weighting process quality. The final weight of each indicator is the final output of this part of the proposal.

3.3.1. Selection of the Weighting Method

In this case, the Fuller's Triangle method see [29] is proposed to be used with a light modification made by the authors of this paper. The modification lies on its comparison scale and allows the possibility of assigning the same importance to two compared attributes and, also, to compute the consistency level achieved during the comparison of attribute pairs. The established scale for the attribute's paired comparison is as follows:

- $a_{ij} = 2$ if the attribute i is considered more important than the attribute j .
- $a_{ij} = 1$ if the attribute i has the same importance as the attribute j .
- $a_{ij} = 0$ if the attribute i is considered to be less important than the attribute j .

3.3.2. Preparation of the Group of Experts for the Weighting

Each expert will be instructed on the characteristics that the attributes will be compared upon. In this case, the relative importance of one indicator over another one will be determined as a result of how this indicator sees and assumes the importance of the PM—MM integration.

3.3.3. Realization of the Paired Comparisons

This step defines the way to perform the comparisons between pairs of attributes. There are two possibilities: (1) to obtain the experts' evaluations in an independent way or (2) to obtain these in a joined fashion. It is important to consider any doubts or questions the experts may have with respect to the content of the indicator given that it may significantly influence the weighting process.

3.3.4. Evaluation of the Consistency Index

In this step, all the possible triplets are generated from the attributes' set, and, on the basis of the rules that are next specified, the inconsistency level and the inconsistency units present in each attribute triplet are subsequently defined.

1. Total inconsistency. It is defined when, in a triplet of attributes, the transitivity principle is totally violated, and this is based on the comparison scale that is proposed, for instance: if it occurs that $a_{ij} = 2$, $a_{jk} = 2$, and $a_{ik} = 0$, in this case, the triplet is assigned or given one inconsistency unit.
2. Partial inconsistency. It is defined when in, a triplet of attributes, the transitivity principle is partially violated, for instance: if it occurs that $a_{ij} = 2$, $a_{jk} = 2$, and $a_{ik} = 1$, in this case, the triplet is assigned or given 0.5 inconsistency units.

3. Consistent. Consistency is assumed when it is not possible to confirm that partial or total inconsistencies are present. In this case, we do not assign inconsistency units to the triplet.

The consistency index (CI) is calculated according to the following equation:

$$CI = 1 - \frac{\sum_i^T U_{ic_i}}{T} \quad (21)$$

where:

U_{ic_i} : units of inconsistency assigned to the attributes' triplet i .

T : number of triplets that are generated by a set of n elements.

$$T = \frac{n!}{3!(n-3)!} \quad (22)$$

With n the number of compared attributes (indicators).

In case the paired comparisons have been performed in an independent way by each expert, then the concordance level reached by these is to be calculated, for which Kendall's concordance coefficient is to be evaluated. In addition, in this case, it is also needed to execute a hypothesis test with respect to the experts' preferences agreement.

3.3.5. Computing of the Weight of Each Attribute

The Simple Ordering method is proposed for use in this step. This method has proved to be an efficient one in this kind of study, see [30]. If the indicators' paired comparisons were performed in an individual way by each expert, the range's mean value would be assumed, and starting from it, a given score will be assigned to each attribute. The authors in [30] back in 2014 presented a methodological tool addressing how to specifically carry out the integration between the Production and Maintenance Management processes at a tactical level. The proposal incorporated elements of the Reliability Centered Maintenance (RCM) philosophy, the Value Analysis, the Ordinal Linguistic Fuzzy Modeling (OFLM), and the Theory of Fuzzy Control. Both the current research and that previous one keep a clear relation although addressing different phases of the whole PM—MM integration.

3.3.6. Analysis of the Results of the Weighting Process

The group of experts will analyze the results and determine if these correspond to the expected differences among the attributes. This will be performed based on the defined feature taken as reference for the weighting itself, i.e., the need for the PM—MM integration.

If there were marked differences, the effectiveness of the steps up until this point will be checked. Otherwise, if results are considered adequate ones, then the methodology moves into the aggregation module, which is presented next.

3.4. Module III: Aggregation of the Attributes

The set of previously evaluated and weighted attributes constitute the input values for this module. The aggregation process, and the obtaining of a main index that defines the level of importance of the production—maintenance integration (PMII)—is the final output of this module.

3.4.1. Definition and Verification of the Optimality Criterion

The optimality criterion will be of the maximum type, which implies that the overall accomplishment of each indicator requires the complete achievement of the concept expressed by the same index PMII. In this sense, the negation operator “Neg” will be applied over the linguistic label that defines the evaluation of those attributes whose raised satisfaction levels imply a low accomplishment of the PMII index.

$$Neg(s_i) = s_j \mid j = g - i \quad (23)$$

3.4.2. Aggregation of the Attributes

The aggregation of the attributes for obtaining the PMII will be realized using the LOWA operator. This will take place using the weights obtained in the second module of the framework. As the application of this operator requires the ordering of the indicators in non-increasing way, in accordance with their weights, in the case of the coexistence of indicators of equal weight values, subsets of these will be created in order to perform partial aggregations of such subsets, generating the weights in accordance with the cardinalities of the terms to aggregate. After this step, the result of these partial aggregations is to be substituted in the initial set of attributes to then perform the global aggregation finally. This strategy we have proposed is different to other approaches in the literature in which the application of the LOWA operator does not consider the possibility of some indicators having equal weight. In addition, in such proposals, the ordering is made arbitrarily, which might produce inconsistent results in many of the cases.

3.4.3. Analysis of the Obtained Result

A qualitative analysis of the obtained result will be carried out at this stage of the methodology in order to determine if it corresponds with expectations. If there are incongruencies in this comparison, the effectiveness of previous steps is to be checked.

4. Results and Discussion

4.1. A Brief Description of the Case Study

The case of study belongs to a manufacturing plant of the Cuban mechanical industry. This plant is dedicated to the production of equipment and spare parts of high priority for the country's economy given that these are mostly used in the sugar, metallurgic, mining, oil, and construction materials industries, to just cite the main ones. The plant produces items such as reducers, rockers, mills, etc., and its manufacturing priorities are directed towards the production cost and delivery time reductions. The plant has a total of 36 machines with a high automatization level.

4.2. Results of the Framework Application

When analyzing the defined set of indicators presented in Table 1, it was decided that the 3rd and 6th ones would be evaluated in a quantitative way, due to the existence of data in the plant, and also the possibility of establishing objective values (goals) for them. The remaining ones were evaluated in a linguistic way. The minimal thresholds to meet in terms of the consensus and consistency indexes were the linguistic term (label) "high" and the value of 0.8, respectively. The results obtained for the indicators evaluated in a qualitative way are the first to be presented next.

Each expert began by assigning each indicator the linguistic label that, in accordance with his/her opinion, represented its current state within the company. The results are shown in Table 2, where the indicators appear numbered in the first column in the exact order they were referred to in Table 1, while the eight selected experts (E1, . . . , E8) appear in the first row. On the other hand, Table 3 shows the result of the main parameters that determine the consensus level (CL_j) achieved in the evaluation with respect to each indicator j , by using the linguistic quantifier "much", with parameters a and b having values of 0.3 and 0.8, respectively.

For the sake of comprehension and demonstrative purposes, an exemplification of the computation of the parameters that determine the consensus level is presented next for the case of the indicator 1. The p_{si1} values are as follows:

$$\begin{aligned} p_{n1} &= p_{mb1} = p_{b1} = p_{m1} = p = \frac{0}{8} = 0, \\ p &= \frac{6}{8} = 0.750, \\ p_{p1} &= \frac{2}{8} = 0.250. \end{aligned}$$

Given that, in this case, according to the used scale, the linguistic terms are consecutive ones, by considering Equation (4) we have that:

$$CR_i = 0.75^2 + 0.25^2 = 0.625$$

Table 2. Evaluations given by the experts to the indicators.

Indicators	E1	E2	E3	E4	E5	E6	E7	E8	F.E
1	vh	p	vh	vh	vh	vh	vh	p	vh
2	h	h	m	h	h	vh	h	h	h
4	vh	vh	h	vh	vh	vh	h	vh	vh
5	h	h	h	h	p	h	h	h	h
7	h	h	h	h	h	h	h	h	h
8	h	h	h	l	m	h	h	h	h
9	m	m	m	l	m	m	m	m	m
10	h	m	h	h	h	h	m	h	h

Table 3. Achieved consensus levels in the evaluation of the qualitative indicators.

Indicators	v_{si}^j	G	p_{nj}	p_{vlij}	p_{lj}	p_{mj}	p_{hj}	p_{vhj}	p_{pj}	CR_j	CL_j
1	{vh, p}	2	0	0	0	0	0	0.750	0.250	0.625	h
2	{m, h, vh}	3	0	0	0	0.125	0.750	0.125	0	0.593	h
4	{h, vh}	2	0	0	0	0	0.250	0.750	0	0.625	h
5	{h, p}	2	0	0	0	0	0.875	0	0.125	0.73	vh
7	{h}	1	0	0	0	0	1	0	0	1	p
8	{l, m, h}	3	0	0	0.125	0.125	0.750	0	0	0.593	h
9	{l, m}	2	0	0	0.125	0.875	0	0	0	0.781	p
10	{m, h}	3	0	0	0	0.250	0.750	0	0	0.625	h

Similarly, in order to obtain CL_1 , Equation (8) was implemented using the linguistic quantifier “much”, whose a and b values were 0.3 and 0.8, respectively.

Before using Equation (10) to obtain the value of l_i , it was necessary to first determine the vector M , which is integrated by those linguistic labels that accomplish or meet the condition defined through Equation (9). In this sense, it was also necessary to determine the membership of the following value with respect to the linguistic terms appearing in Figure 2.

$$\left(\frac{CR_1 - a}{b - a} = 0,65 \right).$$

In this regard, by using Equation (1), it is possible to see that this value presents a membership grade of 0.12 to the linguistic term “medium” (m), 0.88 to the linguistic term “high” (h), and 0 to the remaining terms of the scale. This way, we can conclude that the vector M consists of only one linguistic term, in this case “high”, which coincides with the l_i value and, at the same time, with the consensus level value reached by the experts when evaluating the indicator 1.

As appreciated in Table 3, in all of the cases the obtained consensus level was equal to or higher than the linguistic term “high”. This leads to the conclusion that the demanded quality level for the qualitative attribute evaluation process was met. In the execution of the step 8 of the methodology for obtaining the Global Consensus Level (GLC), we first determined the $\overline{CR}$ through the CR_j values, see Table 3. In this case, we obtained the following value:

$$CR = \frac{\sum_{j=1}^{10} CR_j}{10} = 0.689.$$

By applying the linguistic quantifier Q^2 to this value, in accordance with Equation (12), then we have that the GCL is “very high”, i.e., $GCL = vh$.

It is important to emphasize the fact that our proposal for determining the consensus relation (CR), unlike others, for instance the one presented in [20], has an inclusive character, that is, it considers the opinions of all involved experts and, even more important, it is also sensitive to the proximity among the linguistic terms emitted in evaluation process, as is appreciated when comparing the CR results and the corresponding CL value in the case of the indicators 5 and 9.

In the case of these indicators, it can be seen that there is the same balance in terms of the proportion of experts who provide one or another linguistic label as a measure of the evaluation of both indicators, regardless of whether the linguistic labels used to evaluate are not the same for the indicator. In the case of indicator 5, 87.5% of the experts consider that it is at a “high” level and the 12.5% consider that it is in another state, “perfect”. Similarly, in the case of indicator 9, 87.55% of the experts consider that said indicator is in one state and 12.5% of the members consider that it is presented in another evaluation state. However, since our proposal takes into account the proximity among the linguistic terms (labels), and the labels used in evaluating indicator 9 are closer than those used for indicator 5, then CR_5 , and consequently CL_5 , are also lower than CR_9 and CL_9 respectively.

For both of these indicators the cardinality of the vector v_{si}^j is the same, even the p_{si1} values for the emitted labels in each case are equal, however, given that the proximity among the labels through which indicator 5 was evaluated is inferior, then the CR_5 value is lower than the corresponding one to CR_9 . This fact would not have been detected nor considered if we had simply applied the methodology proposed in [20], in which case the CR value would have been equal to 0.875 for both indicators. This is, again, another of the added values of the presented research.

Another important aspect also lies in the same selection of the proportional linguistic quantifier. Here, it is necessary to find an adequate trade-off between accuracy and application cost, given that the results of CL might be overvalued if the quantifier “At least half” were to be used, or otherwise undervalued if the linguistic quantifier “All” were to be used instead.

The final evaluation of each indicator was obtained using the LOWA operator, as indicated in step 9 of the proposed framework. For the sake of comprehension and demonstrative purposes, an exemplification of the computing or determination is presented next for the case of indicator 1. In Table 3 it is possible to see that the evaluation vector of indicator 1 includes the terms p and vh , with cardinalities of 2 and 6, respectively, which relate to and constitute the values S_6 and S_5 of the used scale’s term set. According to Equation (17), the weights’ vector (w) can be calculated as follows:

$$w_p = \frac{1}{2 \times 2} = 0.25,$$

$$w_{vh} = \frac{1}{2 \times 2} + \frac{1}{2} = 0.75.$$

For the case of two linguistic terms to be aggregated, the adjusted weight vector β_h coincides with the vector W , which is calculated by means of Equation (17). In accordance with Equation (14), the LOWA can be represented as follows:

$$\phi_F(vh, p) = w_1 \otimes s_i \oplus (1 - w_1) \otimes s_j = 0.25 \otimes s_6 \oplus (1 - 0.25) \otimes s_5 = s_k.$$

At the same time, and as mentioned before, this operator requires the non-increasing ordering of the linguistic terms, in accordance with their semantic, and therefore w_1 constitutes the weight associated to the term b_i (the term with the higher semantic value), while $1 - w_1$ represents the weight associated to b_j (the term with the lowest semantic value). By substituting this into Equation (14), k can be determined, and with it, the linguistic term (s_k) resulting from the aggregation.

$$k = \min\{6, 5 + \text{round}(0.25 * (6 - 5))\} = 5,$$

$$s_k = vh.$$

The linguistic values defining the evaluation of the remaining qualitative indicators were also determined in a similar way. These results can be appreciated in the last column of Table 2.

Next, the framework continues with the evaluation of the quantitative indicators, in this case the 3rd and 6th ones. In the case of the 3rd indicator, at the moment of the realization of this study in the company, it was well known that it behaved and reached values of around 90%; however, the defined reference value for it was 100%. As for the 6th indicator, it was also known that 25% of the production resources could be used in equipment repairs, however, the experts defined 30% as the reference value in this case. By implementing Equation (13), the following normalized values were obtained:

$$Vn_3 = 1 - \frac{|1 - 0.90|}{1} = 0.90,$$

$$Vn_6 = 1 - \frac{|0.30 - 0.25|}{0.30} = 0.83.$$

On the other hand, the linguistic quantifier “All” was used for the homogenization process with a and b values equal to 0.5 and 1, respectively. Similarly, by making use of Equations (11) and (10) for the values Vn_3 and Vn_6 , it was possible to determine, in the case of indicator 3, that the M vector included only one element (vh), for that reason, this is the linguistic term that defines its evaluation. In the case of indicator 6, the M vector only included the term h , which also corresponded to its evaluation.

All these results were presented to the experts, and it was concluded that there was an adequate correspondence among the obtained evaluations for each indicator and its current state in the production plant. Subsequently, the application of the framework proceeded with module II. The paired comparisons between indicators were performed by the experts in a joint way, by using the approach described in steps 11 to 13 (see Sections 3.3.1–3.3.3). The results are shown in Figure 3.

As the paired comparisons were performed by the group of experts in a joint way, the weighting process quality was analyzed only by the consistency index. The number of triplets with total inconsistency was five, for instance, the triplet 2-4-8 to just cite one. The number of triplets with partial inconsistency was 21, for instance, the triplet 1-5-7 to just cite another one.

Similarly, making the necessary substitutions in Equation (21), it was also possible to determine the consistency index achieved whose value was 0.87. This value was higher than the minimal established value of 0.8. In Table 4, is also possible to appreciate the weight of each indicator, obtained by the Simple Ordering method.

Table 4. Weight of the indicators.

Indicators	Score	Range	Weight
1	17	10	0.182
2	6	4	0.072
3	4	2	0.036
4	7	5	0.091
5	15	8.5	0.155
6	10	7	0.127
7	8	6	0.109
8	5	3	0.055
9	15	8.5	0.155
10	3	1	0.018

									Indicators
1^2	1^2	1^2	1^1	1^2	1^2	1^2	1^2	1^2	1
2^0	3^0	4^0	5^1	6^0	7^0	8^0	9^0	10^0	
	2^2	2^0	2^0	2^0	2^1	2^2	2^0	2^1	2
	3^0	4^2	5^2	6^2	7^1	8^0	9^2	10^1	
		3^0	3^0	3^0	3^1	3^1	3^0	3^2	3
		4^2	5^2	6^2	7^1	8^1	9^2	10^0	
			4^0	4^1	4^1	4^0	4^0	4^1	4
			5^2	6^1	7^1	8^2	9^2	10^1	
				5^2	5^1	5^2	5^1	5^2	5
				6^0	7^1	8^0	9^1	10^0	
					6^2	6^2	6^0	6^1	6
					7^0	8^0	9^2	10^1	
						7^2	7^0	7^2	7
						8^0	9^2	10^0	
							8^0	8^2	8
							9^2	10^0	
								9^2	9
								10^0	10

Figure 3. Paired comparisons performed by the experts on the set of indicators.

Having, at this point, the evaluation and weight of each indicator, it was time to proceed with the aggregation process, which is detailed in the third module of the framework and leads to the obtaining of the PMII index. Before doing this, it was first necessary to use the negation operator over indicators 9 and 10, given that these were inversely proportional to the PMII index itself. Table 5 summarizes the final evaluation of each indicator, the linguistic term to be aggregated, as well as their weight values.

Table 5. Term to be aggregated and weight of each indicator.

Indicators	Evaluation	Term to Aggregate	Weight
1	vh	vh	0.182
2	h	h	0.072
3	vh	vh	0.036
4	vh	vh	0.091
5	h	h	0.155
6	h	h	0.127
7	h	h	0.109
8	h	h	0.055
9	m	m	0.155
10	h	l	0.018

The n fixed-weight values to aggregate using the LOWA operator (ϕ_F) required the definition of an ordered vector B , where its components are the n linguistic terms ordered in a non-increasing fashion according to their weights. The aggregation process implied the

realization of $n - 1$ iterations, where, in each iteration i , two terms were aggregated, one of these being the result of the aggregation at iteration $i - 1$ and the other one the $(n-i)$ th term of the ordered vector B . The operator ϕ_F was applied over the following B and W vectors, which are:

$$B = (vh, h, m, h, h, vh, h, h, vh, l),$$

$$W = (0.182; 0.155; 0.155; 0.127; 0.109; 0.091; 0.072; 0.055; 0.036; 0.018).$$

According to vector W 's components, one subset of indicators of equal weight was identified. In this case, this subset is formed by indicators 5 and 9. The results of the aggregation of both indicators is shown in Table 6. For these two indicators of equal weight, the components of the vector β_h are equal to 0.5, the same applies to the values w_1 and $1 - w_1$.

Table 6. Aggregation of indicators 5 and 9 through the operator ϕ_F .

Iterations	Terms to Aggregate	i	j	w_1	$1 - w_1$	k	ϕ_F
1	h, m	3	4	0.500	0.500	4	h

If we now use the linguistic term “high” (h) as the evaluation value for the same indicators 5 and 9, the new vector B can be expressed as follows:

$$B = (vh, h, h, h, h, vh, h, h, vh, l)$$

On the other hand, by keeping invariable the weight vector W and using Equation (20), it was possible to generate the corresponding adjusted weight vector β_h . Using this vector, it was then also possible to generate the weight components w_1 and $1 - w_1$ which were to be used in each aggregating iteration.

$$\beta_h = (0.667; 0.652; 0.681; 0.667; 0.667; 0.668; 0.665; 0.671; 0.667; 0.333)$$

Table 7 shows the result of the aggregation in each iteration. The last iteration provides the PMII index result, i.e., PMII = “very high” (vh). This result was contrasted with the opinions of the group of experts, and their agreement with the result was also verified.

Table 7. Aggregation of the indicators through the operator ϕ_F .

Iterations	Terms to Aggregate	i	j	w_1	$1 - w_1$	k	ϕ_F
1	vh, l	2	5	0.667	0.333	4	h
2	h, h	4	4	0.671	0.329	4	h
3	h, h	4	4	0.665	0.335	4	h
4	vh, h	4	5	0.688	0.312	5	vh
5	h, vh	4	5	0.333	0.667	4	h
6	h, h	4	4	0.667	0.333	4	h
7	h, h	4	4	0.681	0.319	4	h
8	h, h	4	4	0.652	0.348	4	h
9	vh, h	4	5	0.667	0.333	5	vh

A Partial Comparative Analysis in the Calculation of the PMII Index

To partially and further demonstrate the feasibility of our proposal, we proceed to determine the same PMII index by applying the LOWA operator in its original version (F_{Q^1}), as it appears in [20,23,28]. In this case, the weights of the terms to be aggregated are calculated from the concept of the relative quantifier (Q^1).

As explained in Section 3.2.8, the application of the LOWA operator requires a non-decreasing ordering of the vector of linguistic terms B of to be aggregated; this ordering will be performed according to its semantics. Thus, considering the evaluation of each indicator that appears in Table 5, we have:

$$B = (vh, vh, vh, h, h, h, h, h, m, l)$$

In this case, the weight vector will be calculated from the application of the proportional fuzzy quantifier Q^1 , expressed in a numerical domain $Q^1 \in [0, 1]$, as shown in [20]. Specifically, $Q^1(r)$ indicates the degree to which a portion r of objects satisfies the concept expressed by the quantifier Q^1 . This degree of satisfaction is calculated using the following equation:

$$Q^1(r) = \begin{cases} 0 & \text{if } r < a \\ \frac{r-a}{b-a} & \text{if } a \leq r \leq b \\ 1 & \text{if } r > b \end{cases} \quad (24)$$

With a and b the minimum and maximum values that define the semantics of the quantifier ($a, b, r \in [0, 1]$).

In terms of the original LOWA version, the weight w_i of each indicator i if is calculated using Equation (25):

$$w_i = Q^1\left(\frac{j}{n}\right) - Q^1\left(\frac{j-1}{n}\right) \quad (25)$$

where j is the position occupied by indicator i within the ordered vector B .

Table 8 shows the evaluations and weights of the set of ordered indicators of the vector B calculated by Equation (25) according to this alternative comparison method, this using the LOWA operator in its original form. The majority proportional quantifier was used as defined in [20] with $a = 0.3$ and $b = 0.8$.

Table 8. Evaluations and weights of the set of ordered indicators according to the original version of the LOWA operator.

Indicators	Evaluation	Term to Aggregate	Weight	β_h
1	vh	vh	0	0
3	vh	vh	0	0
4	vh	vh	0	0
2	h	h	0.2	0.2
5	h	h	0.2	0.2
6	h	h	0.2	0.2
7	h	h	0.2	0.2
8	h	h	0.2	0.2
9	m	m	0	0
10	h	l	0	0

Table 9 shows the iterations of the aggregation of the indicators of the ordered vector B (from right to left as it is performed in the LOWA operator). As can be seen, indicators 9 and 10 are not considered in the aggregation process since both have a weight equal to zero. This is one of the limitations of the original approach and something we have solved with our modifications.

As a result of the application of this LOWA operator in its original formulation, the PMII index is evaluated as “high”. This result differs from the value obtained from our proposal since the original approach does not take into account the real weight of the indicators, but instead the weight is associated to and depends on its the evaluation of the indicator itself. This is something that in practice may be often far from reality, and thus we have proposed its modification.

Table 9. Aggregation of the indicators through the operator F_{Q^1} .

Iterations	Terms to Aggregate	i	j	w_1	$1-w_1$	k	F_{Q^1}
1	h, h	4	4	0.2	0.8	4	h
2	h, h	4	4	0.2	0.8	4	h
3	h, h	4	4	0.2	0.8	4	h
4	h, h	4	4	0.2	0.8	4	h
5	h, vh	4	5	0.0	1.0	4	h
6	h, vh	4	5	0.0	1.0	4	h
7	h, vh	4	5	0.0	1.0	4	h

To sum this partial comparison up, it is important to once more highlight the impact of the aggregation method proposed in this paper over the final result of the PMII index. In this sense, the idea of identifying subsets of indicators with the same weight and performing partial aggregations of these subsets, as was done in the case of indicators 5 and 9 (see Table 6), in order to later substitute the result into a global aggregation which considers all indicators, produces a different result to the one that would have been obtained in the case of simply establishing an arbitrary ordering of the vector B , by placing the linguistic term that defines the evaluation of indicator 9 (in this case m) before the term that defines the evaluation of the indicator 5 (term h), and considering both had equal weights. If this had been performed this way, the result of the PMII index after the global aggregation would have been “high” (h) instead.

In addition, given the non-linear character and non-decreasing monotonic characteristics of the aggregation operator used, it can be also seen how it produces a result that reflects, to a better extent, the state of the indicators of higher relative importance, which is usually a very favorable element from the practical point of view. In this case study, the indicator 1, evaluated as “very high”, concentrates around 20% of the weight.

Finally, based on the case study and the research presented here, it is also possible for the authors of this paper to conclude that the use of different well-known evaluation methods, as for instance, the Likert numeric scales or numeric operators such as the Weighted Average, would have produced a different result, a result influenced by a lack of flexibility of the numeric evaluation methods and the linear character of the mentioned operator. However, as indicated in the framework’s last step, it is still and always up to the experts involved in the study to evaluate the effectiveness of the results achieved.

5. Conclusions

The present research arose from a deep analysis of the states and practice, where it was verified that, although there had been good efforts and published works, none of the existing approaches had presented a clear method specifically designed to assess the importance of PM—MM integration. In this sense, the present paper presented a set of novel indicators and a framework enabling the realization of their evaluation, weighting, and aggregation. Based on this, it was possible to create and propose an index (PMII) that allows us to evaluate the importance of the PM—MM integration in a company. The conception of this index constitutes a useful tool within the decision-making process, as it justifies the level of efforts aimed at improving the level of integration between both processes. The authors firmly believe this all adds practical, methodological, and scientific value to our research and covers some of the gaps identified in the literature analysis.

The inclusion of Ordinal Fuzzy Linguistic Modeling for the evaluation of qualitative attributes made it possible to give an adequate treatment to the vagueness and imprecision that characterizes the evaluative judgment of experts, in addition to constituting a comfortable medium for the experts themselves when expressing their evaluations. The fact of allowing the possibility of characterizing the consensus achieved among experts also increases the power of this technique.

The proposed framework introduced a novel methodology to evaluate the consensus of the experts based on a proposal previously published in [20]. The proposed methodology includes key modifications regarding (1) the calculation of the consensus relationship by including the concept of distance between fuzzy numbers with a triangular membership function, and (2) a new way of generating the adjusted vector of weights (β_h) used in the LOWA operator for the aggregation. Both of these modifications address limitations in the previous approaches, and thus we also believe that this constitutes part of the scientific and methodological value of our research.

On the other hand, the authors of this paper believe that part of the fundamental contributions of the study also lies in the same proposal of the set of original indicators that synthesize the different aspects of the PM—MM integration in the company. These indicators are, at the same time, the base for the calculation of the PMII index proposed.

Similarly, the weighting method proposed within this paper also includes a small modification of the well-known Fuller's Triangle method to obtain the weights. This modification allows quantification of the level of consistency achieved during the paired comparisons, and thereby assessment of the level of quality achieved in this step. It also offers greater flexibility to the method, since it allows assigning of equal importance to each of the two indicators that are evaluated within the pair, something that was not possible with the scale presented by the original method. In addition, the fact of considering and using both qualitative and quantitative indicators together increases the flexibility of the proposed framework with respect to other partial (not equally oriented or as complete as this) proposals in the state of the art and practice, in which the indicators are presented either in a linguistic domain or in a numerical one.

The application of the proposals of this research into practical case study demonstrated its feasibility for real-world use, constituting a new method to highlight the need to improve PM and MM processes in the company and, even more, to achieve adequate coordination between both processes.

Future work will be oriented to information representation of the modeling using two tuples for the evaluation of qualitative indicators, what would lead to an even more efficient reduction of information losses. Future research will also encompass a feasibility analysis in terms of including in the findings in this paper a model based on linguistic hierarchies and unbalanced linguistic sets for the evaluation of certain indicators. Also, new approaches for the ordering of the linguistic terms to aggregate using the LOWA operator will be a subject of future investigation.

Author Contributions: Conceptualization, R.D.C., D.R.D.S. and E.M.D.L.P.M.; methodology, R.D.C., D.R.D.S., E.M.D.L.P.M. and J.P.; validation, R.D.C., D.R.D.S. and C.D.D.T.; formal analysis, R.D.C., D.R.D.S., E.M.D.L.P.M. and J.P.; investigation, R.D.C., D.R.D.S. and E.M.D.L.P.M.; resources, R.D.C., D.R.D.S., E.M.D.L.P.M. and J.P.; writing—original draft preparation, R.D.C., D.R.D.S. and C.D.D.T.; writing—review and editing, R.D.C., D.R.D.S., E.M.D.L.P.M., J.P. and C.D.D.T.; visualization, R.D.C. and D.R.D.S.; supervision, R.D.C., D.R.D.S., E.M.D.L.P.M. and J.P.; project administration, R.D.C. and D.R.D.S.; funding acquisition, D.R.D.S. and J.P. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by the Ministry of Education, Youth and Sports of the Czech Republic and the Faculty of Mechanical Engineering VŠB-TUO, grant number SP2023/088.

Data Availability Statement: Data are contained within the article.

Acknowledgments: This paper has been developed in connection with the Students' Grant Competition project SP2023/088 "Specific Research of Modern Manufacturing Technologies for Sustainable Economy" financed by the Ministry of Education, Youth and Sports of the Czech Republic and the Faculty of Mechanical Engineering VŠB-TUO. This support and funding are both gratefully acknowledged.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Zhao, Y.; He, Y.; Zhou, D.; Zhang, A.; Han, X.; Li, Y.; Wang, W. Functional risk-oriented integrated preventive maintenance considering product quality loss for multistate manufacturing systems. *Int. J. Prod. Res.* **2021**, *59*, 1003–1020. [CrossRef]
2. Khatab, A.; Diallo, C.; Aghezzaf, E.-H.; Venkatadri, U. Integrated production quality and condition-based maintenance optimisation for a stochastically deteriorating manufacturing system. *Int. J. Prod. Res.* **2019**, *57*, 2480–2497. [CrossRef]
3. Schutz, J.; Chelbi, A.; Rezg, N.; Ben Salem, S. Production and maintenance strategies for parallel machines with load transfer in case of failure. *J. Qual. Maint. Eng.* **2019**, *25*, 525–544. [CrossRef]
4. Wakiru, J.M.; Pintelon, L.; Muchiri, P.; Chemweno, P. Integrated maintenance policies for performance improvement of a multi-unit repairable, one product manufacturing system. *Prod. Plan. Control* **2021**, *32*, 347–367. [CrossRef]
5. Liu, Q.; Dong, M.; Chen, F.F. Single-machine-based joint optimization of predictive maintenance planning and production scheduling. *Robot. Comput.-Integr. Manuf.* **2018**, *51*, 238–247. [CrossRef]
6. Ao, Y.; Zhang, H.; Wang, C. Research of Integrated Decision Model for Production Scheduling and Maintenance Planning with Economic Objective. *Comput. Ind. Eng.* **2019**, *137*, 106092. [CrossRef]
7. Weinstein, L.; Cheng, C.-H. Integrating maintenance and production decisions in a hierarchical production planning environment. *Comput. Oper. Res.* **1999**, *26*, 1059–1074. [CrossRef]
8. Kodali, R.; Mishra, R.P.; Anand, G. Justification of world class maintenance systems using analytic hierarchy constant sum method. *J. Qual. Maint. Eng.* **2009**, *15*, 47–77. [CrossRef]
9. Robson, K.; Trimble, R.; MacIntyre, J. Creating and Sustaining a Maintenance Strategy: A Practical Guide. *J. Bus. Adm. Res.* **2013**, *2*, 77–83. [CrossRef]
10. Robson, K.; Trimble, R.; MacIntyre, J. The inhibitors and enablers of maintenance and manufacturing strategy: A cross-case analysis. *Int. J. Syst. Assur. Eng. Manag.* **2014**, *5*, 107–117. [CrossRef]
11. Garg, A.; Deshmukh, S.G. Flexibility in Maintenance: A Framework. *Glob. J. Flex. Syst. Manag.* **2009**, *10*, 21–34. [CrossRef]
12. Naughton, M.D.; Tiernan, P. Individualising maintenance management: A proposed framework and case study. *J. Qual. Maint. Eng.* **2012**, *18*, 267–281. [CrossRef]
13. Royo Sánchez, J.A.; Berges Muro, L.; Lope Domingo, M.A.; Aguilar Martín, J.J.; González Pedraza, R. Production and Maintenance Integrated Resource Planning Management System. In Proceedings of the 17th OrthoGES European Maintenance Congress, Barcelona, Spain, 11–13 May 2004.
14. Jonsson, P. Towards an holistic understanding of disruptions in Operations Management. *J. Oper. Manag.* **2000**, *18*, 701–718. [CrossRef]
15. Swanson, L. An empirical study of the relationship between production technology and maintenance management. *Int. J. Prod. Econ.* **1997**, *53*, 199–207. [CrossRef]
16. Pinjala, S.K.; Pintelon, L.; Vereecke, A. An empirical investigation on the relationship between business and maintenance strategies. *Int. J. Prod. Econ.* **2006**, *104*, 214–229. [CrossRef]
17. McKone, K.E.; Schroeder, R.G.; Cua, K.O. The impact of total productive maintenance practices on manufacturing performance. *J. Oper. Manag.* **2001**, *19*, 39–58. [CrossRef]
18. Crespo Márquez, A.; Gupta, J. Contemporary maintenance management: Process, framework and supporting pillars. *Int. J. Manag. Sci.* **2006**, *34*, 313–326. [CrossRef]
19. Espinilla, M.; Andrés, R.; Martínez, F.J.; Martínez, L. A 360 degree performance appraisal model dealing with heterogeneous information and dependent criteria. *Inf. Sci.* **2013**, *222*, 459–471. [CrossRef]
20. Herrera, F.; Herrera-Viedma, E.; Verdegay, J.L. A model of consensus in group decision making under linguistic assessments. *Fuzzy Sets Syst.* **1996**, *78*, 73–87. [CrossRef]
21. Xu, Z. EOWA and EWG Operators for aggregating linguistic labels based on linguistic preference relations. *Int. J. Uncertain. Fuzziness Knowledge-Based Syst.* **2004**, *12*, 791–810. [CrossRef]
22. Peláez, J.I.; Doña, J.M.; Gómez Ruiz, J.A. Analysis of OWA operators in decision making for modelling the majority concept. *Appl. Math. Comput.* **2007**, *186*, 1263–1275. [CrossRef]
23. Herrera, F.; Herrera-Viedma, E. Aggregation Operators for Linguistic Weighted Information. *IEEE Trans. Syst. Man Cybern. Part A Syst. Hum.* **1997**, *27*, 646–656. [CrossRef]
24. Cabrerizo, F.J.; Alonso, S.; Herrera-Viedma, E. A consensus model for group decisions making problems with unbalance fuzzy linguistic information. *Int. J. Inf. Technol. Decis. Mak.* **2008**, *8*, 109–131. [CrossRef]
25. Wang, L.; Xue, H. Group Decision-Making Method Based on Expert Classification Consensus Information Integration. *Symmetry* **2020**, *12*, 1180. [CrossRef]
26. Labella, Á.; Rodríguez, R.M.; Alzahrani, A.A.; Martínez, L.A. Consensus Model for Extended Comparative Linguistic Expressions with Symbolic Translation. *Mathematics* **2020**, *8*, 2198. [CrossRef]
27. Morillas, A. Introducción al Análisis de Datos Difusos. Curso de Doctorado en Economía Cuantitativa. Departamento de Estadística y Econometría, Universidad de Málaga, Málaga, Spain. 2000. Available online: <https://www.studocu.com/latam/document/universidad-tecnologica-de-la-habana-jose-antonio-echeverria/ingenieria-de-software-i/antonio-morillas-rama-introduccion-al-analisis-de-datos-difusos/17491208> (accessed on 18 December 2023).
28. Delgado, M.; Verdegay, J.L.; Vila, M.A. On Aggregation Operations of Linguistic Labels. *Int. J. Intell. Syst.* **1993**, *8*, 351–370. [CrossRef]

29. Agarski, B.; Budak, I.; Kosec, B.; Hodolic, J. An approach to multi-criteria environmental evaluation with multiple weight assignment. *Environ. Model. Assess.* **2012**, *17*, 255–266. [CrossRef]
30. Díaz Cazañas, R.; Delgado Sobrino, D.R. On the integration of production and maintenance planning at the tactical level: Proposal of a contribution procedure. *Appl. Mech. Mater.* **2014**, *474*, 35–41. [CrossRef]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

Article

A Method to Handle the Missing Values in Multi-Criteria Sorting Problems Based on Dominance Rough Sets

Ahmet Topal ^{1,2,*}, Nilgun Guler Bayazit ² and Yasemen Ucan ²¹ Department of Mathematics Engineering, Istanbul Technical University, Sariyer, 34469 Istanbul, Türkiye² Department of Mathematics Engineering, Yildiz Technical University, Esenler, 34220 Istanbul, Türkiye; guler@yildiz.edu.tr (N.G.B.); ucan@yildiz.edu.tr (Y.U.)

* Correspondence: atopal@itu.edu.tr

Abstract: The handling of missing attribute values remains a challenging and problematic issue in data analysis. Imputation techniques are key procedures used to deal with missing attribute values. However, although these methods are widely used, they cause data bias. Rough set theory, a unique mathematical tool for decision making under uncertainty, overcomes this problem by properly adjusting the relationships. Rough sets are often preferred in both classification and sorting problems. The aim of sorting problems is to sort the objects in the decision table (DT) from best to worst and/or to select the best one. For this purpose, it is necessary to obtain a pairwise comparison table (PCT) from the DT. However, in the presence of missing values, the transformation from DT to PCT is not feasible because there are no ranking methods in the literature for sorting problems based on rough sets. To address this limitation, this paper presents a way to transform from DT to PCT and introduces a generalization of the relation belonging to the “do not care” type of missing values in the dominance-based rough set approach (DRSA) to the decision support tool jRank. We also adapted the DomLem algorithm to enable it to work in PCT with missing values. We applied our method step by step to a decision table with 11 objects and investigated the effect of missing values. The experimental results showed that our proposed approach captures the semantics of ‘do not care’ type missing values.

Keywords: decision tables; pairwise comparison table (PCT); missing attribute values; dominance-based rough sets; ranking

MSC: 03E75; 68V30

1. Introduction

Decision-making under uncertainty is a fundamental challenge in many fields, ranging from economics and business to science and everyday life [1–3]. Unlike decisions made under certainty, where outcomes and probabilities are known, uncertainty requires individuals and organizations to make decisions with incomplete or ambiguous information. This introduces a level of complexity that requires a careful assessment of risks, potential rewards, and unforeseen consequences. Understanding how to approach decisions in uncertain environments is crucial for minimizing risks and maximizing outcomes [1].

In today’s business environment, data are critical elements in making organizations more competitive and developing more effective strategies [4]. As organizations increasingly rely on data-driven insights to guide their operations, the ability to collect, analyze, and interpret data has become a key differentiator [4,5]. When used properly, data enable organizations to make more informed, knowledge-based decisions. In parallel with the increasing reliance on data, the need for effective data management has become more prominent. Organizations need to focus not only on collecting large amounts of data but also on ensuring the quality and integrity of the data they collect.

One of the biggest challenges at this point is that data sets often suffer from incomplete or missing information due to various factors, such as human error, technical glitches, or limitations in data collection tools. These missing values pose significant challenges, particularly in decision-making processes where each criterion plays a crucial role in determining outcomes. These missing values need to be filled in order to apply an appropriate multi criteria decision making technique.

Many techniques have been developed in the literature to make it possible to analyze datasets with missing values [6–8]. One of the simplest and most commonly used approaches is to delete the samples with the missing values. This method, which is usually preferred for datasets with a small number of missing values, results in a loss of information in the datasets, as it also deletes the values of other available attributes of the relevant instances [9]. Another commonly used approach is to fill the missing attribute values with an appropriate value (e.g., mean, median) or by using advanced methods such as machine learning techniques and model-based methods [10]. However, these approaches introduce data bias [11]. Therefore, the use of these techniques in decision making is in some cases undesirable or should be treated with caution.

Rough set theory [12], proposed by Pawlak in 1982, allows for analysis without forcing missing attribute values in data tables. The problem is solved by correctly defining the relationship that forms the information granules [13]. Kryszkiewicz [14] adapted the concepts of dispensable and indispensable attributes, core, decision rules, and reductions from classical rough set theory to incomplete information systems with respect to a tolerance relation. Stefanowski and Tsoukias [15] proposed two extensions of classical rough set theory to deal with the missing values: valued tolerance relations and non-symmetric similarity relations. Wang [16] generalized the classical rough set theory based on the limited tolerance relation, making it suitable for incomplete data tables. On the other hand, adaptations of the dominance-based rough set approach (DRSA) under incomplete information systems were presented in [17,18]. Błaszczyński, Słowiński, and Szlag [18] presented alternative ways of handling missing values in dominance-based rough sets. In their experimental study, Naive Bayes, C4.5, Ripper, VC-DomLEM-mv2, VC-DomLEM-smv, SVM, OLM, and OSDL classifiers were compared in terms of mean absolute error on datasets with the percentages of missing values ranging from 5% to 50%. Experimental results revealed that VC-DomLEM-mv2, SVM, and Naive Bayes were the best-performing classifiers.

Unlike the classification problem in dominance-based or classical rough sets, it may be desirable to rank objects in a decision table and choose the best one among them [19]. However, the most important point is how to rank objects under missing attribute values. A gap in the literature is the absence of rule-based ranking for information systems with “do not care” type missing values. The key idea of this paper is to define a methodology for dealing with a multi-criteria ranking problem based on the DRSA using a decision table that contains “do not care” type of missing values. Our proposed approach consists of three main steps: transformation from decision table to PCT, generalization of dominance relations under incomplete PCT, and adaptation of the DomLem algorithm. Our contributions are detailed as follows.

- We constructed a PCT from an incomplete decision table by defining transformation formulas for both ordinal and cardinal attributes.
- We introduced a generalization of dominance relations to compute the dominating set, dominated set, and set approximations under an incomplete PCT.
- We adapted the DomLem algorithm to deal with “do not care” missing values for the purpose of extracting decision rules from a PCT.

The remaining parts of the paper are structured as follows. In Section 2, we provide some basic mathematical notions of the DRSA and its extension to the “do not care” type missing values. In Section 3, we present how a multi-criteria ranking method based on the dominance rough set approach can be applied to an information system with missing values of the “do not care” type. Furthermore, we establish the relationships concerning the transformation of attributes from the decision table to the pairwise comparison table in the

presence of missing values. Additionally, we address the adaptation of the rule induction algorithm DomLEM to pairwise comparison tables with missing values. In Section 4, we present the experimental design and its outcomes through an example. The last section has our concluding remarks.

2. Dominance-Based Rough Set Approach

DRSA, which was introduced by Greco, Matarazzo, and Słowiński [20] in the early 2000s, is a new version of rough set theory suggested for multi-criteria decision analysis. DRSA, rather than the classical rough set theory, does not require the discretization of continuous condition attributes, and it has other distinctive characteristics such as handling inconsistent elements originating in the dominance principle and taking into account the preference order of the attributes [20]. The limitations that the classical approach had on some issues has been eliminated thanks to these improvements. Basically, DRSA attempts to extract useful information from the ordinal evaluations of objects on preference-ordered attributes (or criteria). In this approach, the decision attribute has a preference-ordered domain, and the attributes are divided into two categories depending on the monotonic relationship with the decision attribute: gain-type attributes and cost-type attributes. A gain-type attribute is one in which the high values in its domain are favored over the low values in its domain in terms of the decision. However, a cost-type attribute is a criterion (i.e., its domain has preference-ordered values) that is not a gain-type. For example, we consider the credit risk (decision attribute) of a person with respect to the value of their assets (condition attribute). Here, low credit risk is preferred to high credit risk. Since the person has a low (high) credit risk when the value of assets is high (low), high values of the condition attribute are at least as good as low values according to the decision. So, it is a gain-type criterion. In the continuation of this section, basic mathematical concepts about the DRSA and the adaptation of the DRSA to the “do not care” kind of missing values will be explained in detail [17,20–23].

2.1. Basic Concepts

A decision table is a simple visual representation that describes objects in terms of various attributes. In this table, objects or elements are placed into rows, while attributes are placed into columns. An object’s value for a particular attribute is entered into the cell where the corresponding row and column intersect. In mathematical terms, a decision table is described by the following quadruple [20]:

- A finite non-empty set of objects (samples or elements) E ;
- A finite set of attributes $X = X_c \cup X_d$, where X_c and X_d are condition and decision attribute sets, respectively;
- The set of values is taken by all attributes in the decision table (denoted by V);
- Information function $f : E \times X \rightarrow V$.

To simplify matters, we assume that $X_d = \{x_d\}$. For $E = \{e_1, e_2, e_3\}$ and $X_c = \{x_{c1}, x_{c2}, x_{c3}\}$, Table 1 illustrates a sample decision table.

Table 1. A sample decision table.

Objects (E)	Condition and Decision Attributes ($X = X_c \cup X_d$)			
	x_{c1}	x_{c2}	x_{c3}	x_d
e_1	$f(e_1, x_{c1})$	$f(e_1, x_{c2})$	$f(e_1, x_{c3})$	$f(e_1, x_d)$
e_2	$f(e_2, x_{c1})$	$f(e_2, x_{c2})$	$f(e_2, x_{c3})$	$f(e_2, x_d)$
e_3	$f(e_3, x_{c1})$	$f(e_3, x_{c2})$	$f(e_3, x_{c3})$	$f(e_3, x_d)$

Let O_x represent an outranking relation on the universe set E according to the attribute x . If for $e_1, e_2 \in E$ and $x \in X$, then the expression $e_1 O_x e_2$ means that object e_1 is at least as good as object e_2 according to the attribute x . Since it has transitive and reflexive properties,

it is a pre-order relation. Also, $e_1 D_Q e_2$ refers to “ e_1 Q -dominates e_2 ” if object e_1 is at least as good as object e_2 for all attributes in any subset Q of the set of condition attributes. Similar to the outranking relation, D_Q is a pre-order relation. However, since every pair of objects in the universe set E cannot be compared with respect to the relation D_Q , it is a partial pre-order relation. The relationship between the outranking relation and dominance relation can be expressed as follows:

$$e_1 D_Q e_2 = \bigwedge_{\forall q \in Q \subseteq X_c} e_1 O_q e_2 \quad (1)$$

where $\wedge$ denotes the logical AND operator.

Using the dominance relation, it becomes viable to identify the set of objects dominated by an element $e_i \in E$. This set is known as the Q -dominated set, denoted by $D_Q^-(e_i)$, and defined as $D_Q^-(e_i) = \{e_j \in E : e_i D_Q e_j\}$. On the other hand, the set of objects dominating object e_i can also be obtained. This set is known as the Q -dominating set, denoted by $D_Q^+(e_i)$, and defined as $D_Q^+(e_i) = \{e_j \in E : e_j D_Q e_i\}$.

The decision attribute x_d forms a partition of the universal set E . Let $K = \{1, 2, \dots, p\}$ be the set of values taken by the decision attribute x_d . It partitions the universal set E into p decision classes, and the definitions of these classes are given as $Cl_n = \{e \in E : f(e, x_d) = n\}$. The upward union of preference-ordered classes, denoted by $Cl_n^{\geq}$, and the downward union of preference-ordered classes, denoted by $Cl_n^{\leq}$, are defined as $\bigcup_{s \geq n} Cl_s$ ($n = 2, \dots, p$) and $\bigcup_{s \leq n} Cl_s$ ($n = 1, \dots, p-1$), respectively. In other words, the set of the upward union of a decision class contains the decision classes that are at least as good as itself, and its downward union contains the decision classes that are at most as good as itself.

With respect to $Q \subseteq X_c$, for $n = 2, \dots, p$, the Q -lower approximation of $Cl_n^{\geq}$, denoted by $\underline{Q}(Cl_n^{\geq})$, and the Q -upper approximation of $Cl_n^{\geq}$, denoted by $\overline{Q}(Cl_n^{\geq})$, are:

$$\underline{Q}(Cl_n^{\geq}) = \{e_i \in E : D_Q^+(e_i) \subseteq Cl_n^{\geq}\}, \quad (2)$$

$$\overline{Q}(Cl_n^{\geq}) = \{e_i \in E : D_Q^-(e_i) \cap Cl_n^{\geq} \neq \emptyset\}. \quad (3)$$

Analogously, the Q -lower approximation and the Q -upper approximation of $Cl_n^{\leq}$ for $n = 1, \dots, p-1$ are defined, respectively, as follows:

$$\underline{Q}(Cl_n^{\leq}) = \{e_i \in E : D_Q^-(e_i) \subseteq Cl_n^{\leq}\}, \quad (4)$$

$$\overline{Q}(Cl_n^{\leq}) = \{e_i \in E : D_Q^+(e_i) \cap Cl_n^{\leq} \neq \emptyset\}. \quad (5)$$

The definitions of the Q -boundary region of $Cl_n^{\geq}$ for $n = 2, \dots, p$ and $Cl_n^{\leq}$ for $n = 1, \dots, p-1$ are

$$BN_Q(Cl_n^{\geq}) = \overline{Q}(Cl_n^{\geq}) - \underline{Q}(Cl_n^{\geq}), \quad (6)$$

$$BN_Q(Cl_n^{\leq}) = \overline{Q}(Cl_n^{\leq}) - \underline{Q}(Cl_n^{\leq}). \quad (7)$$

The lower approximation of a set comprises the objects that definitely belong to that set, whereas the upper approximation of a set includes objects that probably belong to that set. Objects present in the upper approximation but absent in the lower approximation constitute the boundary regions. These objects in the boundary regions are referred to as Q -inconsistent elements, as they exhibit some uncertainties.

2.2. Incomplete Information Systems

The symbol “*” will be used in decision tables to indicate any missing attribute values. We will assume that the value of at least one condition attribute for each element e in the universal set is known. In other words, for all $e \in E$, it is possible to find a condition

attribute x_c in set X_c such that $f(e, x_c) \neq *$. We will also assume that the value of each object on the decision attribute is known.

We now define the dominance relations, D_Q and $\sqsubset_Q$ ($Q \subseteq X_c$), before adapting DRSA to information systems with missing values. Given $e_1, e_2 \in E$ without any missing attribute values, and if $e_2 O_q e_1$ for all $q \in Q$, then the object e_2 dominates the object e_1 and is denoted by $e_2 D_Q e_1$. On the other hand, if $e_1 O_q e_2$ for all $q \in Q$, then the object e_2 is dominated by the object e_1 and is denoted by $e_2 \sqsubset_Q e_1$. Therefore, $e_2 D_Q e_1$ if and only if $e_1 \sqsubset_Q e_2$. However, in the presence of the missing values, the dominance relation may lose some of its properties such as transitivity and a specific kind of symmetry [17].

The following outlines the extension of dominance based rough set approach to handle missing attribute values of the “do not care” kind [23]:

- A subject element $\tilde{e}$ dominates a referent object e (denoted by $\tilde{e} D_Q e$) if and only if $\forall q \in Q, \tilde{e} O_q e$, or $f(\tilde{e}, q) = *$, or $f(e, q) = *$.
- A subject element $\tilde{e}$ is dominated by a referent object e (denoted by $\tilde{e} \sqsubset_Q e$) if and only if $\forall q \in Q, e O_q \tilde{e}$, or $f(\tilde{e}, q) = *$, or $f(e, q) = *$.

Dominance relations need to be redefined in order to handle information systems containing missing values. Therefore, the generalized definitions of rough approximations will be utilized.

For an object $e_i \in E$ and $Q \subseteq X_c$, the Q -dominated set (denoted by $\sqsubset_Q^-(e_i)$) and the Q -dominating set (denoted by $\sqsubset_Q^+(e_i)$) corresponding to the relation $\sqsubset_Q$ as well as $D_Q^+(e_i)$ and $D_Q^-(e_i)$ should also be considered. These sets are called negative and positive dominance cones with respect to $\sqsubset_Q$ and their definitions are [23]:

$$\sqsubset_Q^-(e_i) = \{e_j \in E : e_j \sqsubset_Q e_i\}, \quad (8)$$

$$\sqsubset_Q^+(e_i) = \{e_j \in E : e_i \sqsubset_Q e_j\}. \quad (9)$$

$\sqsubset_Q^-(e_i)$ is the set of objects dominated by e_i , whereas $\sqsubset_Q^+(e_i)$ is the set of objects dominating e_i . The important point to emphasize here is that the equalities $\sqsubset_Q^-(e_i) = D_Q^-(e_i)$ and $\sqsubset_Q^+(e_i) = D_Q^+(e_i)$ are always satisfied for each $e_i \in E$ in decision tables that do not contain missing attribute values.

Now, we present the generalized lower and upper approximations under incomplete information systems to address the missing values [13,23]. The generalized Q -lower and Q -upper approximations of the upward union of the decision classes ($Cl_n^{\geq}, n = 2, \dots, p$) are defined, respectively, as follows:

$$\underline{Q}(Cl_n^{\geq}) = \{e \in E : \sqsubset_Q^+(e) \subseteq Cl_n^{\geq}\}, \quad (10)$$

$$\overline{Q}(Cl_n^{\geq}) = \{e \in E : D_Q^-(e) \cap Cl_n^{\geq} \neq \emptyset\}. \quad (11)$$

Analogously, the generalized definitions of the Q -lower and Q -upper approximations of the downward union of the decision classes ($Cl_n^{\leq}, n = 1, \dots, p-1$) are as follows:

$$\underline{Q}(Cl_n^{\leq}) = \{e \in E : D_Q^-(e) \subseteq Cl_n^{\leq}\}, \quad (12)$$

$$\overline{Q}(Cl_n^{\leq}) = \{e \in E : \sqsubset_Q^+(e) \cap Cl_n^{\leq} \neq \emptyset\}. \quad (13)$$

The definitions of Q -boundary regions of the downward and upward union of the decision classes are the same as Equations (6) and (7), respectively.

Decision rules, which provide a concise description of the decision table, are derived from the objects in rough approximations. A decision rule consists of two parts: a cause clause and a decision clause. Decision rules are articulated using the quantifiers “at least”

and “at most”, as attributes have preference-ordered domains. For instance, suppose we have the following decision rule in consideration:

IF Fever of a patient is at least 38.5 °C and loss of sense of taste is at most 60% THEN the patient has at least moderate degree of carrying the SARS-CoV-2 virus.

There are two elementary conditions in the decision rule. First, the patient’s fever is no less than 38.5 °C, and second, the loss of sense of taste is no more than 60%. The conjunction of these elementary conditions forms the cause clause of the decision rule. In the decision clause of the rule, it is stated that the virus-carrying status of a patient with the above characteristics will belong to a middle or higher-level decision class. In our study, we consider the exact decision rules induced from the objects in the lower approximations. The properties and syntax of these rules are as follows [24,25]:

- **Exact $\text{Dec}_{\geq}$ – rule (Type-1)** is extracted from the objects in $\underline{Q}(Cl_n^{\geq})$. Namely, objects that pertain to the lower approximation of the upward union of decision classes are positive, while all the others are negative.
IF $f(e, x_{c1}) \geq r_1$ and $f(e, x_{c2}) \geq r_2$ and $\dots$ and $f(e, x_{cs}) \geq r_s$ THEN $e \in Cl_n^{\geq}$, where $\{x_{c1}, x_{c2}, \dots, x_{cs}\} \subseteq X_c, (r_1, r_2, \dots, r_s) \in V_{c1} \times V_{c2} \times \dots \times V_{cs}$ and $n = 2, \dots, p$.
- **Exact $\text{Dec}_{\leq}$ – rule (Type-3)** is extracted from the objects in $\underline{Q}(Cl_n^{\leq})$. Namely, objects that pertain to the lower approximation of the downward union of decision classes are positive, while all the others are negative.
IF $f(e, x_{c1}) \leq r_1$ and $f(e, x_{c2}) \leq r_2$ and $\dots$ and $f(e, x_{cs}) \leq r_s$ THEN $e \in Cl_n^{\leq}$, where $\{x_{c1}, x_{c2}, \dots, x_{cs}\} \subseteq X_c, (r_1, r_2, \dots, r_s) \in V_{c1} \times V_{c2} \times \dots \times V_{cs}$ and $n = 1, \dots, p - 1$.

If all elementary conditions of a rule r are satisfied by the object e , then the rule r covers the object e . Furthermore, if the object e satisfies both the cause clause and the decision clause of the rule r , then the rule r is supported by the object e .

The DomLEM algorithm [26] needs to be modified in some ways in order to extract decision rules under information systems with the “do not care” kind of missing attribute values [18]. Elementary conditions must be created from the non-missing attribute values of the objects in the rough approximations. Moreover, when the value of an object on an attribute is missing, that object is covered by any elementary condition created from that attribute. In other words, if $f(e, x_c) = *$ for $e \in E$ and $x_c \in X_c$, then all elementary conditions created from the attribute x_c cover the object e . The remaining portions of the algorithm are unchanged from the original version.

3. Material and Methods

Multi-criteria decision analysis includes selection and sorting problems, as well as classification problems. jMAF [27], Jamm [28], and 4emka [25] are all software tools for classification problems, while jRank [19] is a software tool for selection and sorting problems. These tools are convenient and useful in managing the decision-making process. The purpose of jRank is to rank objects in a decision table from best to worst and/or choose the best one [29]. In this section, we consider selection and sorting problems under incomplete decision tables and propose an adaptation of jRank to handle “do not care” type of the missing attribute values.

The dominance relation stands as the only objective information that can be utilized to compare objects [19]. However, many of the objects in the decision table might not be comparable to each other, as the dominance relation is a partial pre-order relation. To make it possible to compare chosen objects pairwise, this weakness of the dominance relation can be remedied with a domain expert. In this case, a pairwise comparison table (PCT) can be prepared by conducting a comprehensive comparison of the reference objects selected in the original decision table by a decision maker [19]. It should also be emphasized that the PCT is a decision table containing pairs of objects:

- Pairs of objects $\{(e_i, e_j) : (e_i \in A) \wedge (e_j \in A) \wedge (A \subseteq E)\}$ are placed in the rows.

- Derived attributes from the original ones are placed in the columns. In the PCT, $X_{PCT} = X_{PCT}^C \cup X_{PCT}^D$, where X_{PCT}^C and X_{PCT}^D represent the set of condition attributes and the set of decision attributes, respectively.
- $\bar{V}$ denotes the set of all values taken by the attributes in the PCT.
- $g : (A \times A) \times X_{PCT} \rightarrow \bar{V}$ is an information function.

Dealing with ranking problems that involve missing values requires transforming from decision table to PCT. We proposed the following definitions to describe how the PCT is constructed using a decision table containing missing values:

The value of pairs of objects on a cardinal attribute is determined by the difference operation. If the value of objects on the cardinal attribute is not missing, then the difference of these values is taken. However, when at least one of the values of the objects on the cardinal attribute is missing, the object pair's value on that attribute is also missing. This case is valid when the objects in a pair are different from each other. On the other hand, the difference operation consistently yields zero for a pair composed of identical objects.

Definition 1. Let e_i and e_j be any two objects, x_C be any cardinal attribute, and g and f be information functions in the PCT and in an ordinary DT (a DT that contains individual objects), respectively. Then, the transformation on the cardinal attribute for a pair of objects is defined in the following manner:

$$g((e_i, e_j), x_C) = \begin{cases} 0 & \text{if } i = j \\ f(e_i, x_C) - f(e_j, x_C) & \text{if } f(e_i, x_C) \neq * \text{ and } f(e_j, x_C) \neq * \text{ and } i \neq j \\ * & \text{if } f(e_i, x_C) = * \text{ and/or } f(e_j, x_C) = * \text{ and } i \neq j \end{cases} \quad (14)$$

On the other hand, the value of a pair of objects on an ordinal attribute is the ordered pair of their values in the original decision table.

Definition 2. Let e_i and e_j be any two objects, x_O be any ordinal attribute, and g and f be information functions in the PCT and in an ordinary DT (a DT that contains individual objects), respectively. Then, the transformation on the ordinal attribute for a pair of objects is defined in the following manner:

$$g((e_i, e_j), x_O) = \begin{cases} (f(e_i, x_O), f(e_j, x_O)) & \text{if } f(e_i, x_O) \neq * \text{ and } f(e_j, x_O) \neq * \\ (*, f(e_j, x_O)) & \text{if } f(e_i, x_O) = * \text{ and } f(e_j, x_O) \neq * \\ (f(e_i, x_O), *) & \text{if } f(e_i, x_O) \neq * \text{ and } f(e_j, x_O) = * \\ (*, *) & \text{if } f(e_i, x_O) = * \text{ and } f(e_j, x_O) = * \end{cases} \quad (15)$$

The decision class of a pair of objects is related to either relation S (comprehensive outranking) or relation S^c (comprehensive non-outranking), depending on the opinion of the decision maker. Given a pair of objects, these relations indicate whether or not the former is preferred over the latter for the decision maker. For example, $e_1 S e_2$ states that object e_1 is at least as good as e_2 , whereas $e_1 S^c e_2$ states that object e_1 is at most as good as e_2 .

Definition 3. Let e_i and e_j be any two objects, x_D be the decision attribute, and g be an information function in the PCT. Then, the decision class of a pair of objects is defined in the following manner:

$$g((e_i, e_j), x_D) = \begin{cases} S & \text{if } e_i S e_j \\ S^c & \text{if } e_i S^c e_j \end{cases} \quad (16)$$

The set of condition attributes X_{PCT}^C can consist of the union of three subsets: the set of regular attributes $X_{PCT}^{C,R}$ (attributes with no preference-ordered domain), the set of ordinal attributes $X_{PCT}^{C,O}$, and the set of cardinal attributes $X_{PCT}^{C,C}$. In our study, we ignore

the set of regular attributes (i.e., $X_{PCT}^{C,R} = \emptyset$). For $Q \subseteq X_{PCT}^C$, we proposed to generalize the dominance relations in [19] under an incomplete PCT as follows:

With respect to set Q , the pair of objects (e_x, e_y) dominates the pair of objects (e_w, e_z) (denoted by $(e_x, e_y)D_Q(e_w, e_z)$) if and only if:

$$\left[(e_x, e_y)D_{X_{PCT}^{C,C}}(e_w, e_z) \iff \forall q_i \in X_{PCT}^{C,C}, \right. \\ \left. \left((e_x, e_y)O_{q_i}(e_w, e_z) \text{ or } g((e_x, e_y), q_i) = * \text{ or } g((e_w, e_z), q_i) = * \right) \right] \\ \wedge \\ \left[(e_x, e_y)D_{X_{PCT}^{C,O}}(e_w, e_z) \iff \forall q_i \in X_{PCT}^{C,O}, \right. \\ \left. \left(e_x O_{q_i} e_w \text{ or } f(e_x, q_i) = * \text{ or } f(e_w, q_i) = * \right) \wedge \left(e_z O_{q_i} e_y \text{ or } f(e_z, q_i) = * \text{ or } f(e_y, q_i) = * \right) \right].$$

With respect to set Q , the pair of objects (e_x, e_y) is dominated by the pair of objects (e_w, e_z) (denoted by $(e_x, e_y)\sqsubset_Q(e_w, e_z)$) if and only if:

$$\left[(e_x, e_y)\sqsubset_{X_{PCT}^{C,C}}(e_w, e_z) \iff \forall q_i \in X_{PCT}^{C,C}, \right. \\ \left. \left((e_w, e_z)O_{q_i}(e_x, e_y) \text{ or } g((e_x, e_y), q_i) = * \text{ or } g((e_w, e_z), q_i) = * \right) \right] \\ \wedge \\ \left[(e_x, e_y)\sqsubset_{X_{PCT}^{C,O}}(e_w, e_z) \iff \forall q_i \in X_{PCT}^{C,O}, \right. \\ \left. \left(e_w O_{q_i} e_x \text{ or } f(e_w, q_i) = * \text{ or } f(e_x, q_i) = * \right) \wedge \left(e_y O_{q_i} e_z \text{ or } f(e_y, q_i) = * \text{ or } f(e_z, q_i) = * \right) \right],$$

where $\wedge$ denotes the logical AND operator, and $D_{X_{PCT}^{C,C}}$, $\sqsubset_{X_{PCT}^{C,C}}$ and $D_{X_{PCT}^{C,O}}$, $\sqsubset_{X_{PCT}^{C,O}}$ denote dominance with respect to set $X_{PCT}^{C,C}$ and dominance with respect to set $X_{PCT}^{C,O}$, respectively. When the PCT contains no missing attribute values, one can easily observe that $(e_x, e_y)D_Q(e_w, e_z)$ implies $(e_w, e_z)\sqsubset_Q(e_x, e_y)$ or vice versa. However, the dominance relations defined above may lose some properties under an incomplete PCT. Therefore, we will provide the following generalized definition of rough approximation:

Given a pair of objects $(e_x, e_y) \in A \times A$ and a subset of condition attributes $Q \subseteq X_{PCT}^C$, positive and negative dominance cones with respect to relations D_Q and $\sqsubset_Q$ are as follows:

- The set of objects that dominates the pair of objects (e_x, e_y) with respect to relation D_Q is:

$$D_Q^+(e_x, e_y) = \left\{ (e_w, e_z) \in A \times A : (e_w, e_z)D_Q(e_x, e_y) \right\}. \quad (17)$$

- The set of objects that dominates the pair of objects (e_x, e_y) with respect to relation $\sqsubset_Q$ is:

$$\sqsubset_Q^+(e_x, e_y) = \left\{ (e_w, e_z) \in A \times A : (e_x, e_y)\sqsubset_Q(e_w, e_z) \right\}. \quad (18)$$

- The set of objects that is dominated by the pair of objects (e_x, e_y) with respect to relation D_Q is:

$$D_Q^-(e_x, e_y) = \{(e_w, e_z) \in A \times A : (e_x, e_y) D_Q (e_w, e_z)\}. \quad (19)$$

- The set of objects that is dominated by the pair of objects (e_x, e_y) with respect to relation $\mathcal{Q}_Q$ is:

$$\mathcal{Q}_Q^-(e_x, e_y) = \{(e_w, e_z) \in A \times A : (e_w, e_z) \mathcal{Q}_Q (e_x, e_y)\}. \quad (20)$$

It should be noted that $D_Q^+(e_x, e_y) = \mathcal{Q}_Q^+(e_x, e_y)$ and $D_Q^-(e_x, e_y) = \mathcal{Q}_Q^-(e_x, e_y)$ are always valid when the PCT has no missing attribute values.

Recall that there are two classes, S and S^c , in the pairwise comparison table. Since there are no other classes that are at least as good as S except itself, the upward union of the class S is equal to itself only (i.e., $(S)^\geq = S$). Similarly, the downward union of the class S^c is also equal to itself (i.e., $(S^c)^\leq = S^c$) because no other classes exist that are at most as good as it. We shall, therefore, provide the definitions of the lower and upper approximations for the decision classes themselves.

The Q -lower approximation of the outranking relation S , denoted by $\underline{Q}(S)$, and the Q -upper approximation of the outranking relation S , denoted by $\overline{Q}(S)$, are:

$$\underline{Q}(S) = \{(e_x, e_y) \in A \times A : \mathcal{Q}_Q^+(e_x, e_y) \subseteq S\}, \quad (21)$$

$$\overline{Q}(S) = \{(e_x, e_y) \in A \times A : D_Q^-(e_x, e_y) \cap S \neq \emptyset\}. \quad (22)$$

The Q -lower approximation of the non-outranking relation S^c , denoted by $\underline{Q}(S^c)$, and the Q -upper approximation of the non-outranking relation S^c , denoted by $\overline{Q}(S^c)$, are:

$$\underline{Q}(S^c) = \{(e_x, e_y) \in A \times A : D_Q^-(e_x, e_y) \subseteq S^c\}, \quad (23)$$

$$\overline{Q}(S^c) = \{(e_x, e_y) \in A \times A : \mathcal{Q}_Q^+(e_x, e_y) \cap S^c \neq \emptyset\}. \quad (24)$$

The Q -boundary region of the outranking relation S , denoted by $BN_Q(S)$, and the Q -boundary region of the non-outranking relation S^c , denoted by $BN_Q(S^c)$, are as follows:

$$BN_Q(S) = \overline{Q}(S) - \underline{Q}(S), \quad (25)$$

$$BN_Q(S^c) = \overline{Q}(S^c) - \underline{Q}(S^c). \quad (26)$$

Note that when $(e_x, e_y) D_Q (e_w, e_z)$ implies $(e_w, e_z) \mathcal{Q}_Q (e_x, e_y)$ and vice versa, then:

- The lower and upper approximations of the outranking relation S , as defined in [19], are identical to the definitions of the lower and upper approximations provided by (21) and (22).
- The lower and upper approximations of the non-outranking relation S^c , as defined in [19], are identical to the definitions of the lower and upper approximations provided by (23) and (24).

Decision rules can be used to offer a broad description of the preference-ordered information in PCTs. Thus, a PCT can be viewed as the set of decision rules in the form of "IF {elementary condition(s)} THEN {decision(s)}". A decision rule is simply a combination of the elementary condition(s) and the decision(s). There are three kinds of rules that can be induced from rough approximations: exact decision rules, possible decision rules, and approximate decision rules [26]. In our experimental work in the next section, we will employ exact decision rules, whose syntaxes are given below [19].

Decision rules derived from the objects belonging to the lower approximation of the outranking relation (Type-1):

IF $f(e_x, q_{i1}) - f(e_y, q_{i1}) \geq r_1$ and $\dots$ and $f(e_x, q_{ie}) - f(e_y, q_{ie}) \geq r_e$ and $f(e_x, q_{i(e+1)}) \geq r_{i(e+1)}$ and $f(e_y, q_{i(e+1)}) \leq s_{i(e+1)}$ and $\dots$ and $f(e_x, q_{ip}) \geq r_{ip}$ and $f(e_y, q_{ip}) \leq s_{ip}$, THEN $e_x S^c e_y$, where $X_{PCT}^{C,C} = \{q_{i1}, \dots, q_{ie}\} \subseteq Q$, $X_{PCT}^{C,O} = \{q_{i(e+1)}, \dots, q_{ip}\} \subseteq Q$, $(r_1, \dots, r_p) \in \bar{V}_{q_{i1}} \times \dots \times \bar{V}_{q_{ip}}$, and $(s_{i(e+1)}, \dots, s_{ip}) \in \bar{V}_{q_{i(e+1)}} \times \dots \times \bar{V}_{q_{ip}}$.

Decision rules derived from the objects belonging to the lower approximation of the non-outranking relation (Type-3):

IF $f(e_x, q_{i1}) - f(e_y, q_{i1}) \leq r_1$ and $\dots$ and $f(e_x, q_{ie}) - f(e_y, q_{ie}) \leq r_e$ and $f(e_x, q_{i(e+1)}) \leq r_{i(e+1)}$ and $f(e_y, q_{i(e+1)}) \geq s_{i(e+1)}$ and $\dots$ and $f(e_x, q_{ip}) \leq r_{ip}$ and $f(e_y, q_{ip}) \geq s_{ip}$, THEN $e_x S^c e_y$, where $X_{PCT}^{C,C} = \{q_{i1}, \dots, q_{ie}\} \subseteq Q$, $X_{PCT}^{C,O} = \{q_{i(e+1)}, \dots, q_{ip}\} \subseteq Q$, $(r_1, \dots, r_p) \in \bar{V}_{q_{i1}} \times \dots \times \bar{V}_{q_{ip}}$, and $(s_{i(e+1)}, \dots, s_{ip}) \in \bar{V}_{q_{i(e+1)}} \times \dots \times \bar{V}_{q_{ip}}$.

If a pair of objects $(e_x, e_y) \in A \times A$ satisfies all elementary conditions of a rule r , it is covered by the rule r . Moreover, if (e_x, e_y) is covered by the rule r and belongs to the decision class suggested by r , then it supports the rule r .

Our method introduces the necessary changes to the algorithm for extracting decision rules from the PCT with missing attribute values. The changes required in the algorithm are similar to the modified DomLEM algorithm used in [18]. The algorithm needs to be rearranged to construct a candidate elementary conditions set and to determine whether any pairs of objects are covered by elementary conditions of a rule under an incomplete PCT. The set of candidate elementary conditions must consist of non-missing attribute values of related objects. In other words, it is possible to create a candidate elementary condition from attribute q_i of a pair of objects (e_x, e_y) when $f(e_x, q_i) \neq *$ and $f(e_y, q_i) \neq *$ (i.e., $g((e_x, e_y), q_i) \neq *$). On the other hand, a pair of objects meets all elementary conditions on the cardinal attribute q_i if at least one of the values of the objects on q_i is missing. Furthermore, a pair of objects meets all elementary conditions on the ordinal attribute q_i if both of the values of the objects on q_i are missing. When one of the values of the objects on the ordinal attribute q_i is known, the object that has a missing value on q_i satisfies all the elementary conditions created from that attribute. In order to determine whether or not the pair of objects is covered by the elementary conditions created from q_i , the other object that does not have a missing value on q_i must also be tested on these elementary conditions. The rest of the algorithm is the same as the original one.

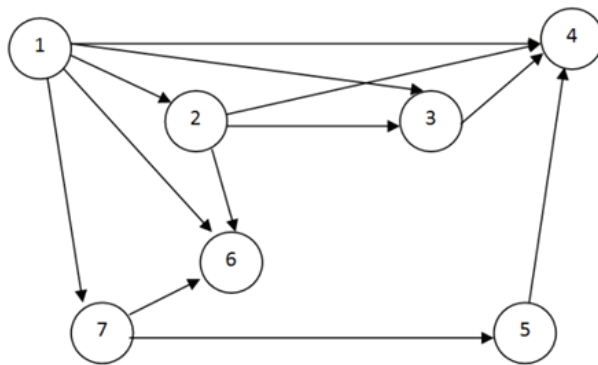
4. Experimental Results and Discussion

In this section, we will consider the incomplete version of the house location problem given in [19]. This example is provided to illustrate the consistency of the generalized relations in our proposed approach under the semantics of ‘do not care’ type of missing values. Table 2 was obtained by deleting some of the attribute values of several locations from the decision table associated with the original house location problem. In this incomplete decision table, both distance and price are cardinal attributes, whereas comfort is an ordinal attribute. The symbols ‘ $\downarrow$ ’ and ‘ $\uparrow$ ’ stand for cost-type attributes and gain-type attributes, respectively. As can be seen from Table 2, distance (q_1) and price (q_2) are cost-type attributes, while comfort (q_3) is a gain-type attribute. Regarding the domain of the attribute comfort, medium is worse than good, but it is better than basic. The opinion of the decision maker for the first seven referent objects (L1–L7) in the decision table is illustrated in Figure 1 in the form of a graph data structure. Each node represents the objects in the house location problem (such as L1, L2, etc.), and the arc between two nodes represents the preference information between the objects related to that arc. For example, node 1 stands for the city Poznan. The arc between node 1 and node 2 means that the city Poznan is at least as good as the city Kapalica (i.e., L1SL2). It will also be assumed that the preference information yielded by the decision maker is symmetric (i.e., L1SL2 $\iff$ L2S^cL1).

Table 2. Incomplete information system for house location problem.

	Distance ($q_1, \downarrow$)	Price ($q_2, \downarrow$)	Comfort ($q_3, \uparrow$)
L1-Poznan	3	60	Good
L2-Kapalica	35	30	Good
L3-Krakow	7	85	Medium
L4-Warszawa	10	90	Basic
L5-Wroclaw	*	60	Medium
L6-Malbork	50	*	Medium
L7-Gdansk	5	70	Medium
L8-Kornik	50	40	Medium
L9-Rogalin	15	50	*
L10-Lublin	*	60	Good
L11-Torun	100	50	Medium

Note: Asterisk (*) indicates the 'do not care' type missing value for the respective attributes.

**Figure 1.** Opinion of the decision maker for the first seven referent objects [19].

The pairwise comparison table in accordance with the preference information shown in Figure 1 is presented in Table 3. For instance, the attribute value for the pair of objects (Poznan, Malbork) in terms of distance was obtained by subtracting the distance values of the former from the latter. Since the value of the location Malbork on the cardinal attribute price is missing, the value of row 4 on this attribute in the PCT is also missing. On the other hand, the value of comfort for the pair of objects (Poznan, Malbork) in the PCT is expressed as an ordered pair. As seen from Figure 1, location Poznan is preferred over the location Malbork by a decision maker so its decision class is the outranking relation S . The Q -dominated sets for the relation D_Q and the Q -dominating sets for the relation C_Q are presented in Appendix A. From Equations (23) and (24), the Q -lower and Q -upper approximations of the non-outranking relation S^c were computed as:

$$\underline{Q}(S^c) = \{13, 14, 15, 16, 17, 18, 19, 20, 21, 22\},$$

$$\overline{Q}(S^c) = \{11, 12, 13, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 27, 29, 30, 31\}.$$

From Equations (21) and (22), the Q -lower and Q -upper approximations of the outranking relation S were computed as:

$$\underline{Q}(S) = \{1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 25, 26, 28\},$$

$$\overline{Q}(S) = \{1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 23, 24, 25, 26, 27, 28, 29, 30, 31\}.$$

By means of the lower and upper approximations of the decision classes in the PCT, the pairs of objects in the Q -boundary regions are:

$$BN_Q(S) = BN_Q(S^c) = \{11, 12, 23, 24, 27, 29, 30, 31\}.$$

Six exact decision rules induced from the Q -lower approximations of S and S^c are listed below. The first three decision rules show the assignments belonging to the decision class S , and the last three rules show the assignments belonging to the decision class S^c .

- IF $f(e_x, q_3) \geq \text{Good}$ and $f(e_y, q_3) \leq \text{Medium}$ THEN $e_x S e_y$.
- IF $f(e_x, q_3) \geq \text{Basic}$ and $f(e_y, q_3) \leq \text{Basic}$ THEN $e_x S e_y$.
- IF $f(e_x, q_1) - f(e_y, q_1) \leq 0$ and $f(e_x, q_3) \geq \text{Good}$ and $f(e_y, q_3) \leq \text{Good}$ THEN $e_x S e_y$.
- IF $f(e_x, q_3) \leq \text{Medium}$ and $f(e_y, q_3) \geq \text{Good}$ THEN $e_x S^c e_y$.
- IF $f(e_x, q_3) \leq \text{Basic}$ and $f(e_y, q_3) \geq \text{Medium}$ THEN $e_x S^c e_y$.
- IF $f(e_x, q_1) - f(e_y, q_1) \geq 32$ and $f(e_x, q_3) \geq \text{Good}$ and $f(e_y, q_3) \leq \text{Good}$ THEN $e_x S^c e_y$.

Considering the assignment part of the rules covered by the pairs of objects, a preference graph was generated using MATLAB and is presented in Figure 2. Nodes in this graph represent all objects in the house location problem, and directed arcs between any two nodes indicate either the relation S or relation S^c . An S -arc between a pair of locations means that this pair of locations is covered by a rule that suggests the assignment to decision class S . Similarly, an S^c -arc between a pair of locations means that this pair of locations is covered by a rule that suggests the assignment to decision class S^c . Directed blue arcs between two locations symbolize the outranking relation S , while directed red arcs symbolize the non-outranking relation S^c . The weight of each directed arc on the graph is also equal to 1.

Table 3. PCT for the house location problem under missing attribute values.

No	Pair	Distance	Price	Comfort	Decision
1	(L1, L2)	−32	30	(Good, Good)	S
2	(L1, L3)	−4	−25	(Good, Medium)	S
3	(L1, L4)	−7	−30	(Good, Basic)	S
4	(L1, L6)	−47	*	(Good, Medium)	S
5	(L1, L7)	−2	−10	(Good, Medium)	S
6	(L2, L3)	28	−55	(Good, Medium)	S
7	(L2, L4)	25	−60	(Good, Basic)	S
8	(L2, L6)	−15	*	(Good, Medium)	S
9	(L3, L4)	−3	−5	(Medium, Basic)	S
10	(L5, L4)	*	−30	(Medium, Basic)	S
11	(L7, L5)	*	10	(Medium, Medium)	S
12	(L7, L6)	−45	*	(Medium, Medium)	S
13	(L2, L1)	32	−30	(Good, Good)	S^c
14	(L3, L1)	4	25	(Medium, Good)	S^c
15	(L4, L1)	7	30	(Basic, Good)	S^c
16	(L6, L1)	47	*	(Medium, Good)	S^c
17	(L7, L1)	2	10	(Medium, Good)	S^c
18	(L3, L2)	−28	55	(Medium, Good)	S^c
19	(L4, L2)	−25	60	(Basic, Good)	S^c
20	(L6, L2)	15	*	(Medium, Good)	S^c

Table 3. Cont.

No	Pair	Distance	Price	Comfort	Decision
21	(L4, L3)	3	5	(Basic, Medium)	S^c
22	(L4, L5)	*	30	(Basic, Medium)	S^c
23	(L5, L7)	*	−10	(Medium, Medium)	S^c
24	(L6, L7)	45	*	(Medium, Medium)	S^c
25	(L1, L1)	0	0	(Good, Good)	S
26	(L2, L2)	0	0	(Good, Good)	S
27	(L3, L3)	0	0	(Medium, Medium)	S
28	(L4, L4)	0	0	(Basic, Basic)	S
29	(L5, L5)	0	0	(Medium, Medium)	S
30	(L6, L6)	0	0	(Medium, Medium)	S
31	(L7, L7)	0	0	(Medium, Medium)	S

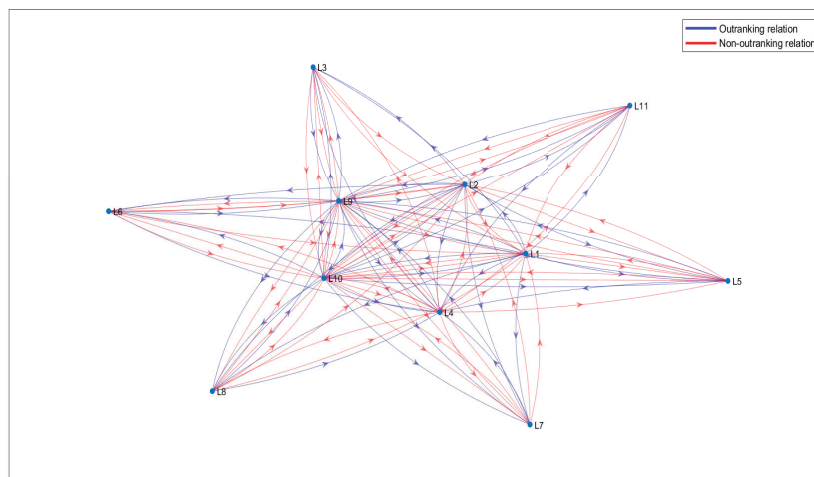


Figure 2. Preference graph for incomplete house location problem.

Six ranking techniques can be applied to the preference graph depicted in Figure 2 to obtain the final rankings of the locations in the decision table [19]. In this experimental study, the Net Flow Score, which is the default technique for the jRank tool, was considered, and the final ranking is provided in Table 4. The Net Flow Score is denoted by $NFS(S, S^c)$ and calculated as follows [19]:

We shall first define the terms positive and negative scores. A positive score refers to the sum of the number of S^c -arcs entering a node and the number of S -arcs leaving that node. A negative score refers to the sum of the number of S^c -arcs leaving a node and the number of S -arcs entering that node. The Net Flow Score for any node is defined as the difference between its positive score and its negative score. Thus,

$$NFS(S, S^c) = PS(S, S^c) - NS(S, S^c), \quad (27)$$

where PS and NS are the functions belonging to the positive and the negative scores, respectively. For example, we consider the location Lublin (L10). As shown in Figure 2, this location has 10 S^c -arcs entering and 10 S -arcs leaving, so its positive score is 20 (i.e., $PS(S, S^c) = 20$). On the other hand, it has 3 S -arcs entering and 10 S^c -arcs leaving, resulting in a negative score of 13 (i.e., $NS(S, S^c) = 13$). Therefore, $NFS(S, S^c) = 20 - 13 = 7$.

Table 4. Final rankings of the objects according to Net Flow Score.

Rank	Locations	$NFS(S, S^c)$
1	L1	17
2	L2	12
3	L10	7
4	L5	−1
5	L9	−2
6	L3, L6, L7, L8, L11	−3
7	L4	−18

As observed from Table 4, locations are sorted in descending order according to their Net Flow Scores. Among the locations, Poznan has the highest score, while Warszawa has the lowest score. Therefore, Poznan and Warszawa are the best and worst locations, respectively. There are five locations with the same Net Flow Score. Line number 6 is shared by the locations Krakow, Malbork, Gdansk, Kornik, and Torun.

By comparing the incomplete decision table to the decision table presented in [19], we can make some important inferences. For instance, in the decision table provided in [19], the value of the location Rogalin for comfort is basic, which is the worst value in the domain of the attribute comfort. However, the value of Rogalin for comfort is missing in Table 2. Since this missing attribute value is of the “do not care” kind, it can take all possible values within the domain of comfort. This missing attribute value can thus have values that are at least as good as basic. It results in the location Rogalin having a better score in this experimental study than its score in [19]. Furthermore, Wroclaw receives almost the best value in the domain of the attribute distance in the decision table provided in [19]. The value of Wroclaw for the attribute distance, however, is missing in Table 2 and this missing value can take any value within the domain of distance. Therefore, the value of this missing attribute could be almost as good as the value it receives in [19]. This leads to the Net Flow Score of Wroclaw being lower than its Net Flow Score computed in [19].

It is clear that a ‘do not care’ type of missing attribute value indicates that it can take any possible value within its domain. The absence of any attribute value for an object directly impacts its final ranking score. In this case, an increase in the number of missing attributes can either positively or negatively influence the final score. The effect of missing attribute values is highly dependent on how these values relate to the overall attribute domain. When the missing values mostly correspond to the worst (or best) values within the defined attribute domain, the object’s score will increase (or decrease) compared to the table without missing values.

Our approach improves the computational complexity cost in various phases, such as rule extraction, dominating set computations, and dominated set computations. It achieves this by reducing the number of comparison operations. For example, in dominating and dominated set computations, when we compare an object with other objects based on any attribute, if the value for that attribute is missing in the object being compared, the object is considered to be dominant over all other objects with respect to that attribute, without any comparison being necessary.

Remark 1. *The efficiency of computational complexity is enhanced through reducing the number of comparisons.*

A strong correlation exists between the final ranking and the consistent and inconsistent pairs of objects in the PCT. The final ranking of consistent objects aligns with the decision maker’s opinion-based ranking. For instance, in the PCT, the object pair (Poznan, Kapalica) is consistent, and as anticipated, Poznan scores higher than Kapalica in the final ranking. On the other hand, (Gdansk, Wroclaw) is an inconsistent pair of objects. In the

final ranking, these two locations must be ranked in a way contrary to the decision maker's opinion. As seen in the table, Wrocław's ranking is at least as good as Gdańsk's. A similar case can also be observed for the other inconsistent pair (Gdańsk, Malbork). This indicates that findings of our methodology aligns with the study [30].

5. Conclusions

Data collected to facilitate the management of decision-making processes may in some cases contain a certain amount of missing values. Although imputation techniques allow data analysis on incomplete datasets, they can lead to data distortion. Rough set theory is able to tackle the data distortion issue by appropriately adapting the relations it employs to construct rough approximations. In the literature, there is no ranking method for information systems with missing attribute values that uses the dominance-based rough set approach. Inspired by this deficiency, in this paper we present how a multi-criteria ranking method based on the dominance-based rough set approach can be applied to an information system with missing values of “do not care” type. Its impact on the final ranking of objects is studied and discussed in detail through an example. Experimental results showed that an object with missing attribute values may experience a change, either an increase or a decrease, in its ranking score compared to when all attribute values are present. It suggests that when the missing values are close to the worst values within the domain of attributes, the score tends to increase, whereas if they are almost the best, the score tends to decrease. Moreover, it indicates that the final ranking supports the decision maker's choices for consistent pairs, but opposes those for inconsistent pairs. Our approach also reduces the number of comparison operations leading to an improvement in computational complexity.

Author Contributions: Conceptualization, A.T.; Methodology, A.T.; Formal Analysis, A.T., N.G.B. and Y.U.; Software, A.T.; Writing—Original Draft Preparation, A.T., N.G.B. and Y.U.; Visualization, A.T.; Writing—Review and Editing, A.T., N.G.B. and Y.U.; Supervision, N.G.B. and Y.U. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Data Availability Statement: The authors confirm that the data that supports the findings of this study are available within the article. Raw data that support the finding of this study are available from the corresponding author, upon reasonable request.

Acknowledgments: This research has been supported by Yildiz Technical University, Scientific Research Coordinatorship under project code FYL-2020-4001.

Conflicts of Interest: The authors declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

Abbreviations

The following abbreviations are used in this manuscript:

DT	Decision Table
PCT	Pairwise Comparison Table
DRSA	Dominance-based Rough Set Approach
SVM	Support Vector Machines
OLM	Ordinal Learning Model
OSDL	Ordinal Stochastic Dominance Learner
VC	Variable Consistency
NFS	Net Flow Score
PS	Positive Score
NS	Negative Score

Appendix A

Appendix A.1. Q -Dominated Sets with Respect to Relation D_Q

Pair #1: 1, 15, 16, 18, 19, 20
 Pair #2: 2, 5, 11, 14, 15, 16, 17, 20, 21, 22, 23, 24, 25, 26, 27, 29, 30, 31
 Pair #3: 2, 3, 5, 9, 10, 11, 13, 14, 15, 16, 17, 20, 21, 22, 23, 24, 25, 26, 27, 28, 29, 30, 31
 Pair #4: 1, 2, 4, 5, 6, 8, 11, 12, 13, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 25, 26, 27, 29, 30, 31
 Pair #5: 5, 11, 14, 15, 16, 17, 20, 21, 22, 23, 24, 25, 26, 27, 29, 30, 31
 Pair #6: 16, 23, 6, 22, 24, 11, 13
 Pair #7: 6, 7, 10, 11, 13, 16, 22, 23, 24
 Pair #8: 2, 5, 6, 8, 11, 13, 14, 15, 16, 17, 20, 21, 22, 23, 24, 25, 26, 27, 29, 30, 31
 Pair #9: 9, 11, 14, 15, 16, 17, 20, 21, 22, 24, 27, 28, 29, 30, 31
 Pair #10: 9, 10, 11, 12, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 27, 28, 29, 30, 31
 Pair #11: 11, 12, 14, 15, 16, 17, 18, 19, 20, 22, 24
 Pair #12: 11, 12, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 27, 29, 30, 31
 Pair #13: 13, 16
 Pair #14: 16, 20, 14, 15
 Pair #15: 15
 Pair #16: 16
 Pair #17: 16, 17, 20, 14, 15
 Pair #18: 16, 18, 19, 20
 Pair #19: 19
 Pair #20: 16, 20
 Pair #21: 15, 21, 22
 Pair #22: 15, 19, 22
 Pair #23: 11, 12, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 27, 29, 30, 31
 Pair #24: 11, 16, 22, 23, 24
 Pair #25: 16, 17, 20, 25, 26, 14, 15
 Pair #26: 16, 17, 20, 25, 26, 14, 15
 Pair #27: 11, 14, 15, 16, 17, 20, 21, 22, 24, 27, 29, 30, 31
 Pair #28: 15, 21, 22, 28
 Pair #29: 11, 14, 15, 16, 17, 20, 21, 22, 24, 27, 29, 30, 31
 Pair #30: 11, 14, 15, 16, 17, 20, 21, 22, 24, 27, 29, 30, 31
 Pair #31: 11, 14, 15, 16, 17, 20, 21, 22, 24, 27, 29, 30, 31

Appendix A.2. Q -Dominating Sets with Respect to Relation Q_Q

Pair #1: 1, 4
 Pair #2: 2, 3, 4, 8
 Pair #3: 3
 Pair #4: 4
 Pair #5: 2, 3, 4, 5, 8
 Pair #6: 4, 6, 7, 8
 Pair #7: 7
 Pair #8: 4, 8
 Pair #9: 3, 9, 10
 Pair #10: 3, 7, 10
 Pair #11: 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 23, 24, 27, 29, 30, 31
 Pair #12: 4, 10, 11, 12, 23
 Pair #13: 3, 4, 6, 7, 8, 13
 Pair #14: 2, 3, 4, 5, 8, 9, 10, 11, 12, 14, 17, 23, 25, 26, 27, 29, 30, 31
 Pair #15: 1, 2, 3, 4, 5, 8, 9, 10, 11, 12, 14, 15, 17, 21, 22, 23, 25, 26, 27, 28, 29, 30, 31
 Pair #16: 1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 16, 17, 18, 20, 23, 24, 25, 26, 27, 29, 30, 31
 Pair #17: 2, 3, 4, 5, 8, 9, 10, 11, 12, 17, 23, 25, 26, 27, 29, 30, 31
 Pair #18: 1, 4, 18, 23, 10, 11, 12
 Pair #19: 1, 4, 10, 11, 12, 18, 19, 22, 23

Pair #20: 1, 2, 3, 4, 5, 8, 9, 10, 11, 12, 14, 17, 18, 20, 23, 25, 26, 27, 29, 30, 31
 Pair #21: 2, 3, 4, 5, 8, 9, 10, 12, 21, 23, 27, 28, 29, 30, 31
 Pair #22: 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 21, 22, 23, 24, 27, 28, 29, 30, 31
 Pair #23: 2, 3, 4, 5, 6, 7, 8, 10, 12, 23, 24
 Pair #24: 2 3 4 5 6 7 8 9 10 11 12 23 24 27 29 30 31
 Pair #25: 2, 3, 4, 5, 8, 25, 26
 Pair #26: 2, 3, 4, 5, 8, 25, 26
 Pair #27: 2, 3, 4, 5, 8, 9, 10, 12, 23, 27, 29, 30, 31
 Pair #28: 3, 9, 10, 28
 Pair #29: 2, 3, 4, 5, 8, 9, 10, 12, 23, 27, 29, 30, 31
 Pair #30: 2, 3, 4, 5, 8, 9, 10, 12, 23, 27, 29, 30, 31
 Pair #31: 2, 3, 4, 5, 8, 9, 10, 12, 23, 27, 29, 30, 31

References

- Mishra, S. Decision-making under risk: Integrating perspectives from biology, economics, and psychology. *Personal. Soc. Psychol. Rev.* **2014**, *18*, 280–307. [CrossRef] [PubMed]
- Zavadskas, E.K.; Turskis, Z. Multiple criteria decision making (MCDM) methods in economics: An overview. *Technol. Econ. Dev. Econ.* **2011**, *17*, 397–427. [CrossRef]
- Von Winterfeldt, D. Bridging the gap between science and decision making. *Proc. Natl. Acad. Sci. USA* **2013**, *110* (Suppl. S3), 14055–14061. [CrossRef] [PubMed]
- Choi, T.M.; Chan, H.K.; Yue, X. Recent development in big data analytics for business operations and risk management. *IEEE Trans. Cybern.* **2016**, *47*, 81–92. [CrossRef] [PubMed]
- Provost, F. *Data Science for Business: What You Need to Know about Data Mining and Data-Analytic Thinking*; O'Reilly Media, Inc.: Sebastopol, CA, USA, 2013; Volume 355.
- Ramezani, R.; Maadi, M.; Khatami, S.M. A novel hybrid intelligent system with missing value imputation for diabetes diagnosis. *Alex. Eng. J.* **2018**, *57*, 1883–1891. [CrossRef]
- Strike, K.; El Emam, K.; Madhavji, N. Software cost estimation with incomplete data. *IEEE Trans. Softw. Eng.* **2001**, *27*, 890–908. [CrossRef]
- Beretta, L.; Santaniello, A. Nearest neighbor imputation algorithms: A critical evaluation. *BMC Med. Inform. Decis. Mak.* **2016**, *16*, 197–208. [CrossRef] [PubMed]
- Lin, W.C.; Tsai, C.F. Missing value imputation: A review and analysis of the literature (2006–2017). *Artif. Intell. Rev.* **2020**, *53*, 1487–1509. [CrossRef]
- Khan, M.A. A Comparative Study on Imputation Techniques: Introducing a Transformer Model for Robust and Efficient Handling of Missing EEG Amplitude Data. *Bioengineering* **2024**, *11*, 740. [CrossRef] [PubMed]
- Vidal-Paz, J.; Rodríguez-Gómez, B.A.; Orosa, J.A. A Comparison of Different Methods for Rainfall Imputation: A Galician Case Study. *Appl. Sci.* **2023**, *13*, 12260. [CrossRef]
- Pawlak, Z. Rough set theory and its applications to data analysis. *Cybern. Syst.* **1998**, *29*, 661–688. [CrossRef]
- Greco, S.; Matarazzo, B.; Slowinski, R. Dealing with missing data in rough set analysis of multi-attribute and multi-criteria decision problems. In *Decision Making: Recent Developments and Worldwide Applications*; Springer: Berlin/Heidelberg, Germany, 2000; pp. 295–316.
- Kryszkiewicz, M. Rough set approach to incomplete information systems. *Inf. Sci.* **1998**, *112*, 39–49. [CrossRef]
- Stefanowski, J.; Tsoukias, A. Incomplete information tables and rough classification. *Comput. Intell.* **2001**, *17*, 545–566. [CrossRef]
- Wang, G. Extension of rough set under incomplete information systems. In *Proceedings of the 2002 IEEE World Congress on Computational Intelligence, 2002 IEEE International Conference on Fuzzy Systems, FUZZ-IEEE'02. Proceedings (Cat. No. 02CH37291)*, Honolulu, HI, USA, 12–17 May 2002; IEEE: Piscataway, NJ, USA, 2002; Volume 2, pp. 1098–1103.
- Szeląg, M.; Błaszczyński, J.; Słowiński, R. Rough set analysis of classification data with missing values. In *Proceedings of the Rough Sets: International Joint Conference, IJCRS 2017, Olsztyn, Poland, 3–7 July 2017; Proceedings, Part I*; Springer: Berlin/Heidelberg, Germany, 2017; pp. 552–565.
- Błaszczyński, J.; Słowiński, R.; Szeląg, M. Induction of ordinal classification rules from incomplete data. In *Proceedings of the Rough Sets and Current Trends in Computing: 8th International Conference, RSCTC 2012, Chengdu, China, 17–20 August 2012; Proceedings 8*; Springer: Berlin/Heidelberg, Germany, 2012; pp. 56–65.
- Szeląg, M.; Słowiński, R.; Greco, S.; Błaszczyński, J.; Wilk, S. jRank—Ranking using Dominance-based Rough Set Approach. *Newsl. Eur. Work. Group Mult. Criteria Decis. Aiding* **2010**, *3*, 13–15.
- Greco, S.; Matarazzo, B.; Slowinski, R. Rough sets theory for multicriteria decision analysis. *Eur. J. Oper. Res.* **2001**, *129*, 1–47. [CrossRef]
- Slowinski, R.; Greco, S.; Matarazzo, B. Rough set and rule-based multicriteria decision aiding. *Pesqui. Oper.* **2012**, *32*, 213–270. [CrossRef]

22. Uçan, Y.; Topal, A.; Bayazit, N.G. A new method for obtaining the inconsistent elements in a decision table based on dominance principle. *Turk. J. Math.* **2020**, *44*, 561–568.
23. Greco, S.; Matarazzo, B.; Slowinski, R. Handling missing values in rough set analysis of multi-attribute and multi-criteria decision problems. In Proceedings of the New Directions in Rough Sets, Data Mining, and Granular-Soft Computing: 7th International Workshop, RSFDGrC'99, Yamaguchi, Japan, 9–11 November 1999; Proceedings 7; Springer: Berlin/Heidelberg, Germany, 1999; pp. 146–157.
24. Greco, S.; Matarazzo, B.; Slowinski, R. A new rough set approach to evaluation of bankruptcy risk. In *Operational Tools in the Management of Financial Risks*; Springer: Berlin/Heidelberg, Germany, 1998; pp. 121–136.
25. Greco, S.; Matarazzo, B.; Slowinski, R. *Multicriteria Classification by Dominance-Based Rough Set Approach*; Politechnika Poznańska: Poznan, Poland, 2000.
26. Greco, S.; Matarazzo, B.; Slowinski, R.; Stefanowski, J. An algorithm for induction of decision rules consistent with the dominance principle. In *Rough Sets and Current Trends in Computing, Proceedings of the Second International Conference, RSCTC 2000, Banff, AB, Canada, 16–19 October 2000*; Revised Papers 2; Springer: Berlin/Heidelberg, Germany, 2001; pp. 304–313.
27. Błaszczyszński, J.; Greco, S.; Matarazzo, B.; Słowiński, R.; Szelag, M. jMAF-Dominance-based rough set data analysis framework. In *Rough Sets and Intelligent Systems—Professor Zdzisław Pawlak in Memoriam: Volume 1*; Springer: Berlin/Heidelberg, Germany, 2013; pp. 185–209.
28. Slowinski, R. The International Summer School on MCDM 2006. Class Note. Kainan University, Taiwan. Software. 2006. Available online: <https://fcds.cs.put.poznan.pl/IDSS/software/jamm.htm> (accessed on 10 June 2024).
29. Alvarez, P.A.; Ishizaka, A.; Martinez, L. Multiple-criteria decision-making sorting methods: A survey. *Expert Syst. Appl.* **2021**, *183*, 115368. [CrossRef]
30. Szelag, M.S. Application of the Dominance-Based Rough Set Approach to Ranking and Similarity-Based Classification Problems. Ph.D Dissertation, Poznań University of Technology, Poznan, Poland, 2015.

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

Article

A Dynamic Trading Approach Based on Walrasian Equilibrium in a Blockchain-Based NFT Framework for Sustainable Waste Management

Ch Sree Kumar ¹, Aayushman Bhaba Padhy ², Akhilendra Pratap Singh ¹ and K. Hemant Kumar Reddy ^{3,*}

¹ Department of Computer Science and Engineering, National Institute of Technology Meghalaya, Shillong 793003, India; p21cs013@nitm.ac.in (C.S.K.); akhilendra.singh@nitm.ac.in (A.P.S.)

² Department of Computer Science and Engineering, NIST Institute of Science and Technology (Autonomous), Berhampur 761008, India; apb121@nist.edu

³ Department of Computer Science and Engineering, VIT-AP University, Amaravati 522237, India

* Correspondence: hemanth.reddy@vitap.ac.in

Abstract: It is becoming harder to manage the growing amounts of waste generated daily at an increasing rate. These problems require an efficient solution that guarantees effectiveness and transparency and maintains trust within the community. To improve the process of traditional waste management, we proposed a unique solution, “GREENLINK”, which uses a combination of blockchain technology with the concept of zero-knowledge proofs (ZKPs), non-fungible tokens (NFTs), and Walrasian equilibrium. Zero-knowledge proofs (cryptographic protocols) are used to verify organizations and prove compliance (e.g., certification, recycling capacity) without disclosing sensitive information. Through an iterative bidding process, the proposed framework employs Walrasian equilibrium, a technique to balance supply and demand, guaranteeing equitable pricing and effective resource distribution among participants. The transactions and waste management activities are securely recorded on an immutable ledger, ensuring accountability, traceability, and transparency. The performance of the proposed model is evaluated. Parameters like average latency, TPS, and memory consumption are calculated using Hyperledger Caliper (a blockchain performance benchmark framework).

Keywords: non-fungible tokens (NFTs); blockchain; zero-knowledge proofs (ZKPs); waste management; decentralized system; Walrasian equilibrium; dynamic bidding; metamask; blockchain; Hyperledger Caliper

MSC: 91A35

1. Introduction

Every day, huge amounts of waste are generated by residential, commercial, and medical sources. Traditionally, the recycling process of this waste material is handled by government organizations, like municipalities, or private organizations. These entities use their resources, such as garbage trucks and employees, to collect waste from various locations, including houses, apartments, medical institutions, industries, etc. However, this process involves a leak of sensitive/personal information of users, such as names, addresses, payment information, and phone numbers. These data are accessible to employees, posing risks such as planned robberies and misuse of waste materials. Loss of materials, unscientific treatment, improper collection of waste, and ethical problems are just a few of the problems that these organizations face. In this case, the user is basically not aware of the

whole process of waste management, does not have control over their data, and is unaware of the authenticity of the organization in charge. Such things erode trust between companies and users. Several organizations provide end-to-end waste management services like collection, segregation, transportation, recycling, composting, incineration, landfilling, hazardous waste treatment, regulatory compliance, and record keeping. Here, all the data generated during these processes are stored centrally. This has drawbacks, such as a single point of failure, scalability limitations, performance bottlenecks, data privacy concerns, etc. Moreover, the verification of waste recycling companies is also an issue. Organizational backdoor entries pose a major risk of cheating customers.

The advantages of using blockchain in waste management are transparency, traceability, efficiency, environmental impact, accountability, etc. The pros of the non-fungible token-based approach to waste management are traceability and transparency (including acting as a digital passport for waste, real-time tracking, preventing illegal dumping), incentivizing recycling and sustainable practices (including rewarding eco-conscious behavior, gamification of recycling, funding sustainable initiatives), ensuring ethical practice, fair trade (fair compensation for recyclers, verifying the origin of recycled materials) and environmental impact tracking, improved public engagement and trust, community-driven initiatives, reduced fraud and corruption, and optimized resource allocation.

To address these challenges, we propose Greenlink—a blockchain-based NFT waste management framework that streamlines waste recycling by connecting verified companies and organizations. The NFT tokens are accessible to all authorized stakeholders, and verified organizations are available to carry out the recycling process efficiently and securely. Dynamic trading refers to a real-time or near-real-time process in which an automated bidding system determines the price for a good, service, or advertisement slot. It is commonly used in online advertising, e-commerce, and auction platforms. Dynamic bidding optimizes pricing to reflect real-time market conditions, increases efficiency and ROI, and enhances competitiveness in the marketplace. The Walrasian equilibrium proposed in our model is a state in a market where supply equals demand, and every participant (buyer or seller) maximizes their utility at given prices. This concept is instrumental in dynamic trading systems for the efficient allocation of resources, real-time price discovery, maximizing platform revenue, stability and predictability, and multi-item auctions. By leveraging Walrasian equilibrium principles, dynamic trading systems enhance fairness, efficiency, and market effectiveness. The Walrasian equilibrium approach aims to find a price vector that balances supply and demand in a market. In this waste bidding model, prices are iteratively adjusted based on excess supply or demand until equilibrium is reached. The verification process is carried out using zero-knowledge proofs, which allows companies to be verified without revealing their private information to the administrator who sets the criteria. Zero-knowledge proofs are an efficient technique to minimize the information transferred between the prover and the verifier in a cryptographic protocol. The prover can prove through a cryptographic commitment scheme to another party, termed the verifier, without revealing how it is proven. Zero-knowledge proofs, or ZKPs, are a technique for one party to cryptographically prove to another that they know something about a piece of information without disclosing the underlying information. A ZKP in the context of blockchain networks may only disclose on-chain that a piece of secret information is legitimate and is known by the prover with high certainty. It operates when the verifier urges the prover to conduct activities that can only be completed correctly if the prover is aware of the fundamental information. If the one who proves is guessing about the outcome of these acts, the verifier's test will almost certainly prove them incorrect with a degree of probability. Zero-knowledge proofs (ZKPs) are defined by key characteristics that make them a powerful, secure, and private verification tool. Accuracy ensures that a

reliable prover can convince an honest verifier that they know the correct input if a claim is true. Stability guarantees that no unethical prover can deceive the verifier into believing they possess the correct information if the claim is false. Zero expertise means that if the claim is true, the verifier learns nothing more from the prover except that the claim is valid. These features make ZKPs a robust method for proving knowledge without revealing additional information. In interactive ZKPs, the prover and verifier communicate back and forth, allowing the verifier to validate the claim, while in non-interactive ZKPs, the prover generates a proof using an algorithm, which is then validated by another algorithm.

Using lightweight protocols like Mina and Polygon zkEVM greatly benefits GREEN-LINK, making the platform energy-efficient. The Mina protocol and Polygon zkEVM use lightweight zk-SNARKS and layer-2 blockchain scaling solutions. The design of the Mina protocol architecture is unique. It has proven to use less memory and CPU because the blockchain size is fixed at 22kb, regardless of the number of transactions. This is made possible by storing images instead of all transactions. This also eliminates the need to install huge storage systems for running the blockchains' nodes and utilizes less power than POW or POS blockchains like Bitcoin. On the other hand, Polygon zkEVM uses the concept of zk-rollups, which follows the approach of winding or rolling multiple transactions into a single form and performing verifications using ZKPs. This approach reduces the computational power and energy costs associated with the transaction process and verification of the data. By performing the transaction computation operation off-chain, which makes the use of smart contracts less and using efficient cryptographic methods, Polygon zkEVM provides high scalability and lower energy consumption options. The performance benchmarks conducted using Hyperledger Caliper ensure that the framework operates efficiently, with low latency, high TPS (transactions per second), and reduced memory consumption translating to less energy use per transaction. By optimizing resource utilization, the framework minimizes its environmental impact during operation.

The proposed framework, Greenlink, includes scalable and efficient features to alleviate the worries about any bottlenecks the system may experience due to its dependence on smart contracts and zero-knowledge proofs using the following features. (A) Use of scalable blockchain infrastructure: For on-chain validation, the proposed framework, Greenlink, uses Polygon zkEVM, which collects data into a single forum and compresses many transactions into a single batch via zk-rollups. This preserves the system's integrity and security while significantly lowering the strain on the main blockchain. Thanks to the use of zk-rollups, Greenlink can handle thousands of transactions per second (TPS), which makes it capable of handling large user traffic and performing operations even during heavy workloads. (B) The Mina protocol's lightweight architecture: The 22 KB size of the Mina blockchain makes it highly efficient in handling the application's memory usage compared to any other blockchain. The architecture guarantees that the system maintains effectiveness even with increased user demand. (C) Optimization of smart contracts: To reduce computing complexity and gas costs, the proposed framework's smart contracts are designed to only use essential information required for a successful waste management process. The framework minimizes the risk of network congestion and maximizes the execution time by using the latest technologies. (D) Event-driven execution: Due to their event-driven nature, smart contracts save needless computing costs by only operating when certain criteria are satisfied. (E) Enhancements in zero-knowledge proofs (ZKPs): The proposed framework, Greenlink, makes sure that—even with high user loads—the ZKP process does not create a bottleneck by optimizing the timings for ZKP production and verification. Our approach aggregates and verifies ZKPs for several transactions together, significantly lowering the computational cost of the verification process. (F) Layer-2 integration: By enabling off-chain processing for most processes, layer-2 solutions, such as

Polygon zkEVM, lessen dependency on the layer-1 blockchain and prevent congestion. (G) Performance benchmarking: Using Hyperledger Caliper, Greenlink tests metrics, including average latency, TPS, and memory usage, under simulated high loads to assess its performance. This guarantees that the system is highly scalable and offers information about any bottlenecks for future enhancements.

The organization of the paper is as follows: Section 1 provides the introduction; Section 2 includes a review of the literature; Section 3 discusses blockchain-enabled NFT frameworks with zero-Knowledge proof (ZKP)-based verification for waste management; Section 4 discusses their implementation and the results; and Section 5 provides the conclusion.

2. Literature Review

This section reviews the latest work in different areas of non-fungible tokens and works, including OpenZeppelin in blockchain technology, zero-knowledge proofs (ZKPs), and the Mina protocol, which relates to the topic of waste management in some way, either directly or indirectly. Researchers worldwide have used these concepts in various areas such as waste management, supply chain management, healthcare, drug traceability, etc. This section briefly describes the research, diving into broader areas such as trading approach, traceability systems, NFT-based artworks, security, and auction methodology. The role of OpenZeppelin contracts is to define the NFTs' buy, sell, and minting functionalities.

ZKP and blockchain technology have the potential to completely transform the waste management industry and the way the entire process is managed. In [1], Khiem et al. proposed an innovative solution for medical waste management, integrating blockchain, smart contracts, and NFTs. In [2,3], Bulkowska et al. and Jiang et al. discussed and implemented blockchain in waste management. In [4], Tian et al. discussed post-carbon footprint analysis in the art market, and in [5], Mielcarek et al. showcased water nutrient management. In [6], Chiacchio et al. discussed NFTs in the pharmaceutical sector, while in [7], Chiquito et al. discussed waste management in aquaculture, and in [8], Rayes et al. proposed a plastic waste management model. Triet et al. proposed a framework in [9] that combined blockchain with circular economy principles to improve construction waste management. In [10], Ahmed et al. developed a framework to examine Blockchain's potential to enhance waste management in smart cities and discussed security. In [11], Yuan et al. proposed an IoT-based solid waste management system for developing countries, utilizing blockchain-enabled VANETS, UHF technology, and geo-fencing for real-time tracking and communication. Steenmans et al., in [12], examined blockchain's role in plastic waste management, identifying key areas like payment, recycling rewards, waste tracking, and smart contracts.

Blockchain technology offers promising benefits to supply chain management by enhancing transparency, security, and efficiency through its decentralization feature and immutable ledger technology, which eliminates reliance on third-party dependencies. In [13], Farsi et al. examined blockchain-based supply chain systems, highlighting their benefits and key requirements while identifying potential cyber threats and security challenges. Soori et al., in [14], explored how integrating blockchain with IIoT in Industry 4.0 enhances sustainable supply chain management, offering real-time tracking and reduced intermediaries. Aslam et al. proposed a framework for evaluating whether a complex supply chain should adopt blockchain technology in [15]. They empirically analyzed the supply chain practices of Pakistan's oil industry and their impact on operational performance. In [16,17], Dietrich et al. and Duan et al. reviewed the most recent publications on blockchain in supply chain management, classifying them by complexity. They found that most blockchain projects focus on simple supply chains, with no examples yet addressing transparency and

audibility in complex manufacturing supply chains. In [18,19], Dudczyk et al. and Marin et al. discussed blockchain platforms in global supply chain management and distributed technology, their industry applications, and existing solutions.

The adoption of blockchain technology in emerging trading and auctions seems promising and can help enhance efficiency, reduce scams, and form trust between trading parties. Mariia Rodinko et al., in [20], proposed upgrading existing cryptocurrencies to new tokens without relying on external information sources to avoid KYC, ensuring a predicted supply and fair pricing based on decentralized auction simulations. In [21], Cui et al. proposed a novel competition padding auction mechanism for peer-to-peer energy trading among renewable prosumers, addressing the budget deficit issue. In [22], Kadadha et al. proposed AB Crowd, a fully decentralized crowdsourcing framework using the Ethereum blockchain and the repeated single-minded bidder auction mechanism. In [23], Agarwal et al. explored a decentralized blockchain-based double auction mechanism for energy trading in a smart microgrid, allowing prosumers to sell surplus energy to consumers. In [24], Liu et al. addressed computation offloading in mobile blockchain networks by proposing a smart contract-based double auction mechanism.

In [25], Yang et al. proposed a blockchain-based digital identity management system using smart contracts, ZKPs, and a BZDIMS prototype for enhanced privacy. Wali et al., in [26], discussed secure vehicle data management using blockchain and ZKPs. In [27], Soewito et al. study combined AES and ZKPs for authentication, improving data transmission security and reducing authentication processing time. In [28], Diro et al. discussed identity sharing with ZKPs and blockchain for secure identity sharing, addressing vulnerabilities. In [29], Kuznetsov et al. proposed an innovative aggregation scheme for ZKPs within Merkle trees to enhance data verification efficiency in blockchain systems. Experimental results show significant reductions in proof size and computational requirements, offering a scalable and secure solution for various applications. In [30], Xiao Xu et al. explored using blockchain-powered zero-knowledge proofs to create a privacy-preserving disability management system that enhances inclusivity.

The current paper proposes a model that integrates technologies like zero-knowledge proofs (ZKPs), blockchain, and the Mina protocol (a lightweight blockchain that only stores the image of a blockchain) to efficiently handle waste management using a novel dynamic or repeated bidding using Walrasian equilibrium for iterative price adjustments to sell waste and allot private organizations that are verified and genuine. After the ZKP-based verification method, private organizations can participate in the auction process to purchase waste from individuals or organizations in the NFT marketplace. Digital assets are stored via decentralized storage, i.e., NFT.storage, which is a specialized service that provides decentralized functionalities for non-fungible tokens. We have used its API endpoint to successfully store the tokens in its storage. The Coinbase Wallet SDK enables users to connect to Greenlink through the Coinbase Wallet. The Web3Modal library makes connecting decentralized applications to different cryptocurrency wallets simple. The benchmarking tool—Hyperledger Caliper is used to calculate the throughput and efficiency of the entire platform (Table 1).

Table 1. Comparison of the proposed system with other existing systems.

References	Framework/Model Used	Performance Evaluation	Objective
[1]	Hyperledger fabric, NFT framework	Transaction throughput: 200 TPS	Address waste management using blockchain in Vietnam.
[9]	Ethereum-based smart contracts	Reduced waste mis management by 30%	Streamlining medical waste management in healthcare.
[10]	Multi-chain architecture for smart cities	Cost efficiency improved by 20%	Blockchain for smart city waste management.
[11]	Decentralized information management system	Operational efficiency +25%	Sustainable construction and demolition waste management.
[12]	Public blockchain with smart contract governance	Recycling accuracy +15%	Plastic waste management governance using blockchain.
[13]	Blockchain-based supply chain model	Improved security by 40%	Securing supply chains with blockchain.
[14]	Industrial IoT model integrated with blockchain	Reduced energy consumption in SC	Sustainable supply chain management in Industry 4.0.
[15]	Private blockchain adoption model	Reduction in delays by 20%	Blockchain in oil industry supply chains.
[17]	Blockchain maturity model	Increased ROI by 35%	Application of blockchain in SCM.
[18]	Consortium blockchain	Lowered logistics cost by 15%	Survey of blockchain applications in global SCM.
[20]	Proof-of-Burn consensus	System efficiency: 92%	Decentralized cryptocurrency upgrade mechanism.
[21]	Consortium and double auction	P2P energy trading success rate: 95%	Blockchain for energy trading with double auction.
[22]	Auction mechanism integrated with blockchain	Computational efficiency +20%	Crowdsourcing auction mechanism using blockchain.
[23]	Double auction mechanism	Energy savings: +15%	Smart microgrid energy trading using blockchain.
[24]	Smart contract-based computation offloading	Reduced latency by 20%	Mobile blockchain computation offloading.
[25]	Zero-knowledge-proof-based identity scheme	Improved privacy by 30%	Blockchain identity management using ZKP.
[26]	ZKAV framework	Enhanced AV performance	Blockchain for autonomous vehicles.
[27]	Modified ZKP algorithm	Faster authentication by 15%	IoT security system with ZKP based verification.

3. Blockchain-Enabled NFT Framework

Waste management is a process and a set of actions that are required to manage waste, starting from its generation to its final disposal. Effective waste management helps reduce waste generation, increase recycling and reuse, minimize environmental impacts (e.g., pollution, climate change), protect public health and safety, and conserve natural resources. NFTs are paving the way to effectively help in waste management in the following ways: tracking and verification, waste ownership and responsibility, incentivizing recycling, supply chain optimization, carbon credit verification, digital twins for waste management, education and awareness, and decentralized waste management. The NFT

blockchain-based waste management trading framework uses decentralized technology to improve waste buying, selling, and ownership. This framework ensures the organization's verification using the zero-knowledge proof (ZKP) method. Provides secure, transparent transactions, reduces the risk of fraud, and enables seamless ownership transfers using smart contracts. The immutable nature of blockchain records each transaction, providing verifiable provenance for NFTs. The framework also provides an auction-based mechanism for purchasing goods by verified organizations, and financial transactions are stored securely on the blockchain.

3.1. Proposed Mathematical Model

To model a bidding process in which organizations and users purchase waste from companies, we can create a mathematical model involving the following components (Table 2):

Table 2. Notations and descriptions.

Notation	Description
$i \in \{1, \dots, N\}$	Index for waste-selling companies.
$j \in \{1, \dots, M\}$	Index for waste-buying organizations/users.
Q_i	Quantity of waste available from company i (in units).
Q_{ij}	Quantity of waste allocated by company i to organization j (decision variable).
P_{ij}	Bid price per unit of waste submitted by organization j for company i .
C_i	Reserve price per unit set by company i (minimum acceptable price).
B_j	Budget of organization/user j .
$U_j(Q_{ij})$	Utility or value derived by organization j for acquiring Q_{ij} units of waste.
$R_i(Q_i)$	Revenue function for company i .

Assumption 1. 1. The allocation of waste quantity cannot exceed the available supply or budgetary limits.

2. Each organization maximizes its utility while each company maximizes its revenue.
3. Prices are determined based on bidding values and reserve prices.

Objective Functions:

1. Company Objective: Maximize total revenue:

$$R_i = \sum_{j=1}^M P_{ij} Q_{ij}, \quad \text{subject to } P_{ij} \geq C_i, \quad \forall j \quad (1)$$

2. Organization/User Objective: Maximize total utility:

$$U_j = \sum_{i=1}^N U_j(Q_{ij}) - \sum_{i=1}^N P_{ij} Q_{ij}, \quad \text{subject to } \sum_{i=1}^N P_{ij} Q_{ij} \leq B_j, \quad \forall j \quad (2)$$

The constraints of the mathematical model are as follows.

1. Quantity Constraint for Companies:

$$\sum_{j=1}^M Q_{ij} \leq Q_i, \quad \forall i \quad (3)$$

2. Budget Constraint for Organizations/Users:

$$\sum_{i=1}^N P_{ij} Q_{ij} \leq B_j, \quad \forall j \quad (4)$$

3. Non-Negativity:

$$Q_{ij} \geq 0, \quad \forall i, j \quad (5)$$

4. Bid Price Constraint:

$$P_{ij} \geq C_i, \quad \forall i, j \quad (6)$$

Assuming that the problem involves dynamic or repeated bidding, iteratively adjust prices P_{ij} and allocations Q_{ij} to clear the market (i.e., supply equals demand). The Walrasian equilibrium approach seeks a price vector that balances supply and demand in a market. We iteratively adjust prices based on excess demand or supply for this waste bidding model until equilibrium is achieved.

3.2. Mathematical Modelling of the Dynamic Bidding Using Walrasian Equilibrium

1. Initialization: Set an initial price vector

$$P^0 = [P_{ij}^0] \quad (7)$$

where P_{ij}^0 is the price for each pair for company-buyer (i, j) .

Ensure

$$P_{ij}^0 \geq C_i \quad (8)$$

for all i, j .

2. Excess Demand Calculation: At each iteration t , compute the excess demand for each company i based on the current price P^t :

$$Z_i^t = \sum_{j=1}^M D_{ij}(P_{ij}^t) - Q_i \quad (9)$$

where $D_{ij}(P_{ij}^t)$ is the demand from buyer j for waste from company i at price P_{ij}^t , computed as follows:

$$D_{ij}(P_{ij}^t) = \arg \max_{Q_{ij}} (U_j(Q_{ij}) - P_{ij}^t Q_{ij}) \quad (10)$$

subject to $Q_{ij} \leq \frac{B_j}{P_{ij}^t}$.

3. Price Adjustment: Update prices iteratively based on excess demand Z_i^t :

$$P_{ij}^{t+1} = P_{ij}^t + \alpha \cdot Z_i^t \quad (11)$$

where $\alpha > 0$ is a step size controlling the rate of price adjustment.

4. Stopping Criterion: Continue iterations until

$$|Z_i^t| \leq \epsilon, \quad \forall i \quad (12)$$

where ϵ is a small threshold indicating equilibrium (i.e., supply equals demand).

5. Final allocation: Once equilibrium prices P_{ij}^* are found, allocate waste quantities:

$$Q_{ij}^* = D_{ij}(P_{ij}^*) \quad (13)$$

At equilibrium: Prices: equilibrium prices P_{ij}^* and Quantities Q_{ij}^* satisfy

$$\sum_j Q_{ij}^* = Q_i \quad (14)$$

(Supply equals demand for each company).

3.3. Proposed Framework/Approach—Greenlink

We have proposed a blockchain-based approach for waste management called Greenlink. Greenlink exclusively facilitates waste recycling among verified companies and organizations. Using the zero-knowledge proof method; we ensure the authenticity of the participants. Recyclable waste is auctioned off to interested organizations. The highest bidder purchases the waste, which is then sold through our marketplace. The proposed framework uses the Mina protocol, a lightweight blockchain, as it stores only the image of the blockchain. The architecture of the proposed model is shown in Figure 1.

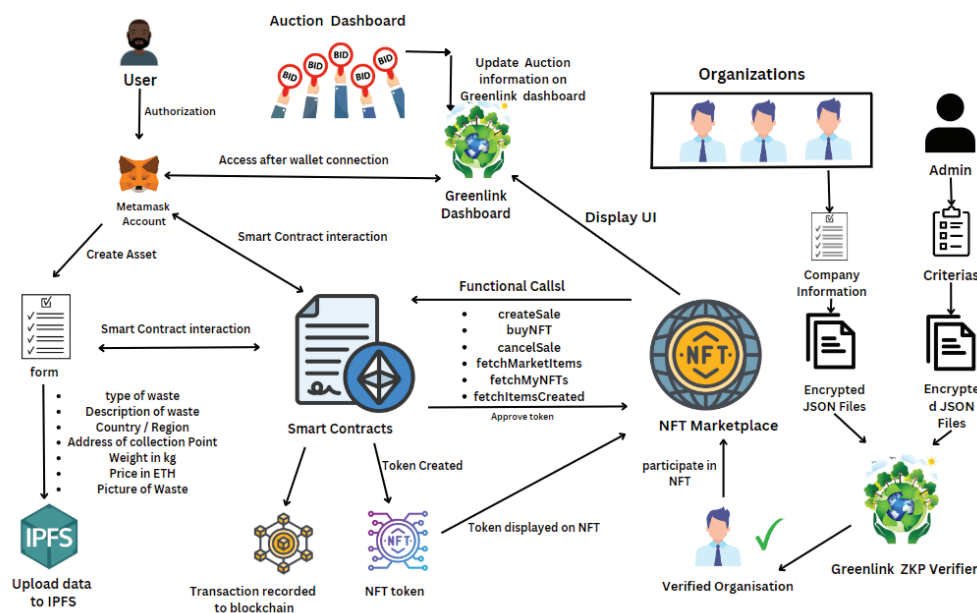


Figure 1. Architecture of the proposed model—Greenlink.

The functionality of the proposed framework—Greenlink—is divided into two parts. The primary functionality is to help people recycle their waste safely and securely through the use of the decentralized marketplace, where users can upload detailed information about waste, such as category, weight, address, etc. The other functionality is to let only valid waste management organizations participate in the decentralized marketplace. Organizations are verified using the zero-knowledge proof (ZKP) method. The proposed marketplace facilitates the sale of recyclable waste with auctions, ensuring competitive pricing and secure transactions for verified participants. The NFT tokens are accessible to all authorized stakeholders, and verified organizations will be available to carry out the recycling process efficiently and securely. Polygon zkEVM provides a scalable environment in which the EVM compatibility and zk-rollups are grouped for the efficient running of smart contracts with Ethereum. A cryptographic primitive known as zero-knowledge proof

(ZKP) is used by the Polygon zkEVM for validating the state transitions. It scales up Ethereum. Transaction confirmations are quicker, and the fees are lower when compared to Ethereum MainNet, which helps design complex blockchain applications.

3.4. Proposed Framework's Algorithms

In our currently developed model, NFT token trading operations were implemented using smart contracts. The following algorithms were developed and used in our framework.

Algorithm 1 shows the minting operation of waste materials and its listing in NFT marketplace.

Algorithm 2 shows the creation of a market item.

Algorithm 3 reflects the transfer of ownership of items and funds between parties.

Algorithm 4 depicts the fetching of unsold market items.

Algorithm 5 shows the fetching operation of items the user has purchased.

Algorithm 6 shows the process of zero-knowledge proof.

Algorithm 7 shows dynamic or repetitive bidding using the Walrasian equilibrium approach.

Algorithm 1 Minting waste and listing in marketplace

```

1: function CREATE_TOKEN(tokenURI (String), price (Integer)) Integer
2:   Increment the token ID counter:
3:   _tokenIds.increment()
4:   Get the current token ID and assign it to newTokenId:
5:   newTokenId = _tokenIds.current()
6:   Mint a new token to the sender with the new token ID:
7:   _mint(msg.sender, newTokenId)
8:   Set the token URI for the new token ID to the provided tokenURI:
9:   _setTokenURI(newTokenId, tokenURI)
10:  Create a market item with the new token ID and the provided price:
11:  createMarketItem(newTokenId, price)
12:  Return the new token ID:
13:  return newTokenId
14: end function

```

Algorithm 2 Creating a market item

```

1: function CREATE_MARKET_ITEM(tokenId: Integer, price: Integer)
2:   if price < 0 then
3:     Raise error "Price must be at least 0 wei"
4:   end if
5:   newMarketItem ← Create MarketItem with:
6:   tokenId = tokenId
7:   seller = sender's address
8:   owner = contract's address
9:   price = price
10:  sold = false
11:  idToMarketItem[tokenId] ← newMarketItem
12: end function

```

Algorithm 3 Transferring ownership of the item, as well as funds between parties

```

1: function CREATEMARKETSALE(tokenId)
2:    $price \leftarrow idToMarketItem[tokenId].price$ 
3:    $seller \leftarrow idToMarketItem[tokenId].seller$ 
4:   if msg.value  $\neq$  price then
5:     Raise error "Please submit the asking price to complete the purchase."
6:   end if
7:    $idToMarketItem[tokenId].owner \leftarrow \text{sender's address}$ 
8:    $idToMarketItem[tokenId].recycled \leftarrow \text{true}$ 
9:    $idToMarketItem[tokenId].seller \leftarrow \text{null address}$ 
10:  Increment the itemsRecycled counter
11:  Transfer the token from contract's address to sender's address
12:  Transfer the payment to the seller
13: end function

```

Algorithm 4 Fetching all unsold market items

```

1: function FETCHMARKETITEMS
2:   return array of MarketItem
3:    $itemCount \leftarrow \text{current value of } \_tokenIds$ 
4:    $unsoldItemCount \leftarrow \text{current value of } \_tokenIds - \text{current value of } \_itemsRecycled$ 
5:    $currentIndex \leftarrow 0$ 
6:   Create a new array items of type MarketItem with length unsoldItemCount
7:   for  $i$  from 0 to  $itemCount - 1$  do
8:     if owner of  $idToMarketItem[i + 1]$  is equal to contract's address then
9:        $currentId \leftarrow i + 1$ 
10:       $currentItem \leftarrow idToMarketItem[currentId]$ 
11:       $items[currentIndex] \leftarrow currentItem$ 
12:       $currentIndex \leftarrow currentIndex + 1$ 
13:    end if
14:  end for
15:  return items
16: end function

```

Algorithm 5 Fetching items that a user has purchased

```

1: function FETCHMYNFTS
2:   Returns array of MarketItem
3:    $totalItemCount \leftarrow \text{current value of } \_tokenIds$ 
4:    $itemCount \leftarrow 0$ 
5:    $currentIndex \leftarrow 0$ 
6:   for  $i \leftarrow 0$  to  $totalItemCount - 1$  do
7:     if owner of  $idToMarketItem[i + 1]$  is equal to sender's address then
8:        $itemCount \leftarrow itemCount + 1$ 
9:     end if
10:  end for
11:  Create a new array items of type MarketItem with length itemCount
12:  for  $i \leftarrow 0$  to  $totalItemCount - 1$  do
13:    if owner of  $idToMarketItem[i + 1]$  is equal to sender's address then
14:       $currentId \leftarrow i + 1$ 
15:       $currentItem \leftarrow idToMarketItem[currentId]$ 
16:       $items[currentIndex] \leftarrow currentItem$ 
17:       $currentIndex \leftarrow currentIndex + 1$ 
18:    end if
19:  end for
20:  return items
21: end function

```

Algorithm 6 Zero-knowledge proof

```

1: Struct Report: OrgID, validUntil, RedeableAmt, hasCondA, hasCondB, hasCondC
2: Struct Requirements: OrgID, verifyTime, minRedeableAmt, maxRedeableAmt, allowCondA, allowCondB,
   allowCondC
3:
4: function HASHREPORT(report, reqToCheck)
5:   return hash(report.OrgID, report.validUntil, report.RedeableAmt,
6:     report.hasCondA, report.hasCondB, report.hasCondC,
7:     reqToCheck.OrgID, reqToCheck.verifyTime,
8:     reqToCheck.minRedeableAmt, reqToCheck.maxRedeableAmt,
9:     reqToCheck.allowCondA, reqToCheck.allowCondB,
10:    reqToCheck.allowCondC)
11: end function
12:
13: procedure VERIFYORGANIZATION(report, reqToCheck)
14:   Step 1: Verify OrgID
15:   Assert report.OrgID == reqToCheck.OrgID
16:   Step 2: Verify validUntil
17:   Assert currentTime ≤ report.validUntil
18:   Assert reqToCheck.verifyTime ≤ report.validUntil
19:   Step 3: Verify RedeableAmt
20:   Assert report.RedeableAmt ≥ reqToCheck.minRedeableAmt
21:   Assert report.RedeableAmt ≤ reqToCheck.maxRedeableAmt
22:   Step 4: Verify Conditions
23:   if reqToCheck.allowCondA then
24:     Assert report.hasCondA == TRUE
25:   end if
26:   if reqToCheck.allowCondB then
27:     Assert report.hasCondB == TRUE
28:   end if
29:   if reqToCheck.allowCondC then
30:     Assert report.hasCondC == TRUE
31:   end if
32:   Step 5: Generate Verification Hash
33:   verificationHash = HASHREPORT(report, reqToCheck)
34:   Step 6: Store Verification Status
35:   State.verifiedReqHash[verificationHash] = TRUE
36:   return verificationHash
37: end procedure

```

The description of the algorithms developed in our framework is given below:

Algorithm 1 is about minting waste and listing it in the marketplace. It creates a token and returns the token ID and the price. Based on this, the createMarketItem function is called upon, as discussed in Algorithm 2.

Algorithm 2 creates a market item with the following information: tokenID, sender's address, contract's address, and price.

Algorithm 3 transfers ownership of the item as well as funds between the parties. This function, named createmarketSale, stores the sender's address and transfers the token from the contract's address to the sender's address. It transfers the payment to the seller.

Algorithm 4 fetches the unsold marketItems function and keeps track of the unsold items in the marketplace. Here, if the owner of the marketItem is equal to the contract's address, then the counter is increased, and the item is fetched based on the ID.

Algorithm 5 fetches the items that the user has purchased. Here, if the owner of the marketItem is verified with the sender's address, then the item is fetched.

Algorithm 6 shows the zero-knowledge proof (ZKP) algorithm.

Algorithm 7 Dynamic or repetitive bidding using Walrasian equilibrium.**function** AUCTION Input: N = Number of companies and M = Number of buyers $Q_i = [Q_1, Q_2, \dots, Q_N]$ = Supply for each company $B_j = [B_1, B_2, \dots, B_M]$ = Budgets for each buyer $C_i = [C_1, C_2, \dots, C_N]$ = Reserve prices for each company $U_j(Q_{ij})$ = Utility function for buyer j α = Step size for price adjustment and ϵ = Convergence threshold

max_iterations = Maximum number of iterations

Output: P_{ij} = Final price matrix and Q_{ij} = Final allocation matrix**Steps:**

Step 1: Initialization

Initialize price matrix P_{ij} with reserve prices C_i :

$$P_{ij}^0 = C_i \quad \text{for all } i, j$$

Initialize iteration counter: $t = 0$ Initialize excess demand $Z_i^t = [0, 0, \dots, 0]$ (size N)

Step 2: Iterative Adjustment

Repeat until convergence or maximum iterations:

1. Compute demand $D_{ij}(P_{ij}^t)$ for each buyer j and company i :

$$D_{ij} = \arg \max_{Q_{ij}} (U_j(Q_{ij}) - P_{ij}^t \cdot Q_{ij})$$

$$\text{Subject to: } Q_{ij} \leq \frac{B_j}{P_{ij}^t} \quad (\text{Budget constraint})$$

2. Calculate excess demand Z_i^t for each company i :

$$Z_i^t = \sum_{j=1}^M D_{ij} - Q_i$$

3. Update prices P_{ij}^t for all i, j :

$$P_{ij}^{(t+1)} = P_{ij}^t + \alpha \cdot Z_i^t$$

4. Check convergence:

$$\text{If } |Z_i^t| \leq \epsilon \quad \text{for all } i : \quad \text{Break the loop}$$

5. Increment iteration counter: $t = t + 1$

Step 3: Final Allocation

For each $i \in [1, N]$:For each $j \in [1, M]$:

$$Q_{ij} = D_{ij}(P_{ij}^t)$$

Step 4: Output Results Return final price matrix P_{ij} and allocation matrix Q_{ij} **end function**

The verification process is carried out using the process of zero-knowledge proofs, which allows companies to be verified without revealing their private information to administrators, who set the criteria. Zero-knowledge proofs are an efficient technique to minimize the amount of information transferred between the prover and the verifier in a cryptographic protocol. The ZKP algorithm has a struct report with the following

fields: OrgID (organization ID), valid until (expiry or validation time), RedeableAmt (redeemable amount), hasCondA, hasCondB, hasCondC (conditions or flags for additional constraints). The constraints are as follows: OrgID (organization ID to check against), verifyTime (time of verification), minRedeableAmt and maxRedeableAmt (limits on the redeemable amount), and allowCondA, allowCondB, allowCondC (allowed conditions for flexibility in validation).

Functions: hashReport(report: Report): Computes a hash value based on the contents of a report. The hash serves as an immutable, verifiable representation of the report.

VERIFYORGANIZATION (report: Report, reqToCheck: Requirements): Validates whether a report satisfies a set of requirements using a zero-knowledge proof approach.

Validates the report using the requirements by:

1. Asserting conditions like matching OrgID (OrgID field in the report matches the OrgID specified in the requirements (reqToCheck)), validUntil time constraints (ensures that the report is not expired and is valid for the requested verification period), and redeemable amount ranges (ensures that only reports with an acceptable amount of redeemable value are validated).
2. Three conditional checks verify whether the specific conditions in the report align with the allowable conditions in the requirements (hasCondA, hasCondB, hasCondC). This ensures that any extra conditions required for validation (e.g., specific certifications, flags) are satisfied.
3. Computing a final verification hash of both the report and the requirements. The hash acts as a unique proof of verification, encapsulating all the validated constraints without exposing the raw data. This is critical for privacy in a zero-knowledge proof system.
4. Storing the hash in State.verifiedReqHash to mark the request as verified.

In Algorithm 7, the verified organizations participate in the auction process. In our framework, a Walrasian equilibrium is used for iterative price adjustments in dynamic bidding within an auction-based waste trading system. The algorithm is given here.

4. Implementation and Result Discussion

4.1. Implementation

The utilization of zkEVM blockchain technology has enabled the process of waste management and the involvement of private organizations in recycling the waste generated, while prioritizing security measures. The proposed architectural framework is developed through the use of zkEVM, Node.js, Ether.js, Web3Modal, Wallet Connect SDK, and NFT.storage API endpoints. The development of decentralized applications requires the implementation of Ethereum, a secure, open-source, scalable, and programmable blockchain platform that offers smart contract development functionalities. The designed and implemented Greenlink framework represents the process of user-generated waste management by verified companies through NFTs. It also adds a layer of security by keeping users, administrators (who set criteria), and organizational identities private. This framework has the potential to introduce a new way of waste management, where trust and privacy are backed up by auction-based NFT trading.

The integration of the various tools and technologies used and their integration within our proposed framework is listed here:

1. Ether.js: Interacts with the Ethereum blockchain. It also provides functionalities to manage wallets, handle transactions, etc.
2. zkEVM: The proposed architectural framework is developed using zkEVM. It facilitates cost-effectiveness and scalability in blockchain transactions and enhances

waste management operations. Quicker transaction processing is facilitated by the duo zero-knowledge rollups (zk-rollups) and the Ethereum Virtual Machine (EVM), ensuring lower fees and that security is also not compromised. The zk-rollups further help in waste management transaction confirmation. The zkEVM and EVM efficiently handle dynamic bidding and NFT issuance for waste-tracking operations.

3. Wallet Connect allows decentralized applications to connect and interact with cryptocurrency wallets securely.
4. Web3Modal is a library that makes connecting decentralized applications to different cryptocurrency wallets simple.
5. NFT.storage: The decentralized storage functionalities for NFT metadata are provided by the NFT's highly dependable storage service. It also helps in tamper-proof record storage for waste management (like origin and recycling history). This enables the stakeholders to confirm the waste sources (i.e., traceability) and the process of recycling. NFT.storage is a specialized service designed to provide decentralized storage functionalities for non-fungible tokens. This also makes for very cost-effective data management.
6. Coinbase Wallet SDK enables users to connect the Coinbase Wallet to Greenlink.
7. Hyperledger Caliper is a performance evaluation benchmarking tool.

For deployment on the blockchain, the address of the current model's contract owner is used. The private key is given in checksum form. To access and use it, the user can mint the token. The Ethereum nodes hold the private key. In the caliper, the signed transactions and the data are sent to the node when the private key is available for the address of the model's contract owner. The address of the benchmark invokes the methods in the contract. The `network_config` file accesses the `from_address` and its `private_key` with a password to invoke the methods in the caliper benchmark.

The caliper adapter hosts the successful transactions over the network along with the block count. The JSON object is given by the `config_list` to be deployed on the network prior to executing the caliper benchmark. JSON entries for the `createMarketItem`, `createMarketSale`, `createToken`, `fetchItemListed` smart contracts of our model are specified, and this enables the invocation of the methods on the contracts.

4.2. Result Discussion

The proposed NFT marketplace model's performance is evaluated using the Hyperledger Caliper, and the Polygon blockchain is used over Ethereum network. The OS is Ubuntu 18.04. The various tools are listed in Table 3.

Performance metrics like the speed of transactions, tokenization (which ensures precision and reliability), and consumption of resources (CPU and memory usage) are monitored. These are evaluated based on the transaction rates for the network's scalability and efficiency measurement. This helps us to evaluate and assess the real-world performance of our proposed model.

The benchmarking of each step is performed by providing an input size (load) of around 5000 transactions. Table 4 shows the overall performance of the proposed model by Hyperledger Caliper, which reflects the fact that every transaction was successfully performed over the Ethereum network. This table shows the send rate, latency, and throughput of the proposed model's four essential smart contracts.

Table 3. List of tools/software used with their version.

Sl#	Requirements	Specifications
1	React	v17.0.0
2	Vite	v2.9.9
3	Ethers.js	v5.6.6
4	Web3Modal	v1.9.7
5	Coinbase Wallet SDK	v3.1.0
6	Solidity	v0.8.7
7	NFT.storage	v6.2.0
8	Wallet Connect	v1.7.8
9	NPM	v10.5.0
10	Vs Code	v1.92.0

The current model's performance is tested with a total transaction count of 5000. The input transaction rates are 1000, 2000, 3000, 4000, and 5000 per second. The CreateMarketItem, createMarketSale, createToken, and fetchItemListed functions' performance is noted. The number of workers varied from 1–5. A total of 15 evaluation runs were conducted (i.e., for workers 1–5 for every transaction rate, i.e., 1000, 2000, 3000, 4000, and 5000 per second). The performance data are summarized in Tables 4 and 5. The various performance metrics like latency, throughput, memory usage, CPU usage, traffic in/out, and disk read/write are plotted in the graphs in the results subsection. The utilization of resources like CPU, memory, throughput, and traffic flowing in and out are shown in Table 5.

Table 4. Send rate, average latency, and throughput of the four functions used in our model.

Function Name	Txns	Success	Send Rate (TPS)	Avg Latency (s)	Throughput (TPS)
create Market Item	5000	5000	1000	8.76	248.8
	5000	5000	2000	8.75	249.9
	5000	5000	3000	8.72	252.8
	5000	5000	4000	8.71	253.4
	5000	5000	5000	8.74	254.2
create Market Sale	5000	5000	1000	10.37	238.8
	5000	5000	2000	9.96	239.5
	5000	5000	3000	9.84	241.8
	5000	5000	4000	10.4	242.6
	5000	5000	5000	10.34	243.1
create Token	5000	5000	1000	11.05	230.7
	5000	5000	2000	11.09	230
	5000	5000	3000	10.89	230
	5000	5000	4000	11.05	230.5
	5000	5000	5000	11.02	230.6
fetch Item Listed	5000	5000	1000	10.51	227.6
	5000	5000	2000	10.48	227.1
	5000	5000	3000	10.07	228.4
	5000	5000	4000	10.43	227.4
	5000	5000	5000	10.51	227.9

Table 5. CPU utilization, memory utilization, and traffic in/out of the four functions with respect to the send rate 1000 to 5000 TPS.

Function Name	Send Rate (TPS)	CPU% (max)	CPU% (avg)	Memory (max) [MB]	Memory (avg) [MB]	Traffic In [MB]	Traffic Out [MB]
create Market Item	1000	93.9512	51.18625	791	745	12.9	28.2
	2000	92.9511	50.5484	825	772	12.8	26.9
	3000	91.2542	49.2547	888	780	13.1	28.9
	4000	90.9087	47.2548	902	826	12.9	27.8
	5000	93.2588	47.5457	988	872	12.9	28.1
create Market Sale	1000	25.7462	4.3125	989	784	9.95	31.2
	2000	26.5871	4.4154	992	810	9.81	30.2
	3000	27.5844	4.7871	993	809	9.73	30.2
	4000	28.5857	4.7985	987	842	9.93	31.2
	5000	29.9899	4.8215	991	859	10.02	31.1
create Token	1000	28.0651	4.6914	22.065	4.69	12.1	32.5
	2000	28.4362	4.7678	23.854	4.76	12.1	32.6
	3000	28.8981	4.8125	23.754	5.01	11.9	31.6
	4000	29.8785	4.9254	23.896	4.81	12.1	32.5
	5000	29.8025	4.9891	24.021	4.91	12.2	32.4
fetch Item Listed	1000	23.2541	4.1254	22.854	4.6535	10.3	31.5
	2000	24.8765	4.3658	23.459	4.1625	10.2	31.4
	3000	26.5898	4.4989	23.898	4.7965	10.3	32.2
	4000	27.0054	4.5175	23.878	4.5175	10.2	31.3
	5000	28.4575	4.7585	23.587	4.7585	10.3	31.4

Table 5 represents the analysis of the latency for operations like createMarketItem, createMarketSale, createToken, and fetchItemListed, respectively. These functions were measured at 1000, 2000, 3000, 4000, and 5000 transactions per second (TPS), respectively. The latency details in Table 4 are plotted in Figure 2a, which reveals that average latency is relatively high for the CreateToken function (deep blue color), whereas it is lower for the createMarketItem function (light blue color). Another trend also reveals that as the send rate (transactions per second) increases, the average latency gradually decreases. The average latency reaches 8.71 s when the workers are set at 3. The average latency reaches a maximum of 8.76 s when the no. of workers = 4 for the CreateMarketItem function. Similarly, the average latency ranges from a maximum of 10.37 s to a minimum of 9.84 s. The createToken function's average latency range ranges from a minimum of 10.89 s to a maximum of 11.05 s. For the function fetchItemListed, the maximum average latency is 10.51 s and the minimum is 10.07 s.

The graph in Figure 2b shows the throughput for createMarketItem, createMarketSale, createToken, and fetchItemListed, respectively, when the transaction per second is varied from 1000 to 5000, as shown in Table 4. In Figure 2b, the createToken function has a lower throughput compared to the other three tokens (deep red color in Figure), and the createMarketItem has the highest throughput (light blue color) of all the functions for all the rates of transactions (1000 to 5000). The graphs in Figure 2a,b have been drawn with different color codes to differentiate the four functions, and they are grouped together for each of the transaction rates 1000–5000 TPS, respectively.

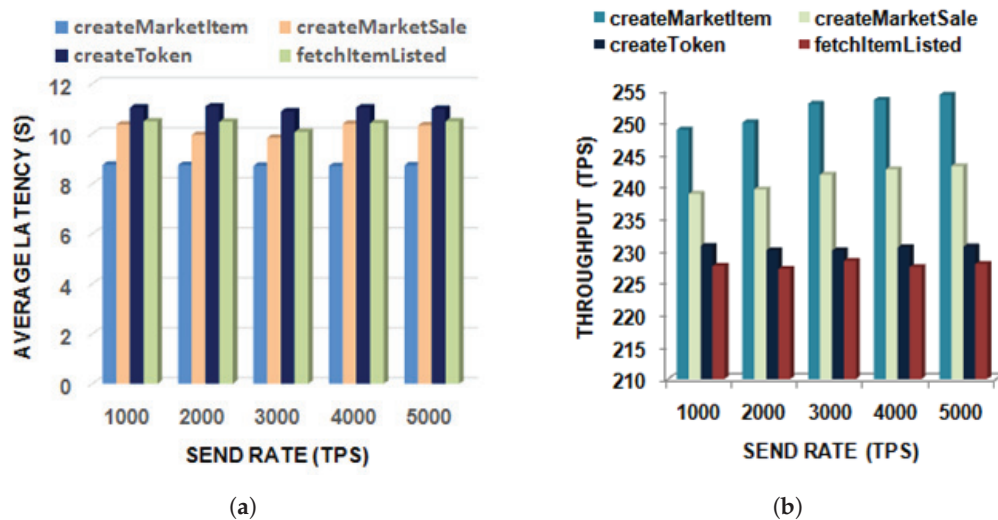


Figure 2. (a) Average latency; (b) throughput. (a) shows the average latency of the solidity functions createMarketItem, createMarketSale, createToken, and fetchItemListed; w.r.t send rate (in TPS). (b) shows the throughput of the solidity functions createMarketItem, createMarketSale, createToken, and fetchItemListed; w.r.t send rate (in TPS).

Figure 3a shows the graph of the memory utilization in MB versus the send rate (1000–5000 TPS) of the createMarketItem function. The graph shows that the maximum memory utilization is 791 MB when the send rate is 1000 TPS, and it is 988 MB when the send rate is 5000 TPS. For the createMarketSale function, Figure 3b shows the maximum memory utilization is 989 MB when the send rate is 1000 TPS, and it is 991 MB when the send rate is 5000 TPS. For the createToken function, as shown in Figure 3c, the maximum memory utilization is 22.0 MB when the send rate is 1000 TPS, and it is at 24.0 MB when the send rate is 5000 TPS. In the fetchItemListed function, the graph in Figure 3d shows the maximum memory utilization is 22.8 MB when the send rate is 1000 TPS, and it is 23.5 MB when the send rate is 5000 TPS. The maximum memory consumption for the createToken and fetchItemListed functions is low for all the transaction rates in the range of 1000–5000 TPS (22.0 MB to 24.0 MB for both functions).

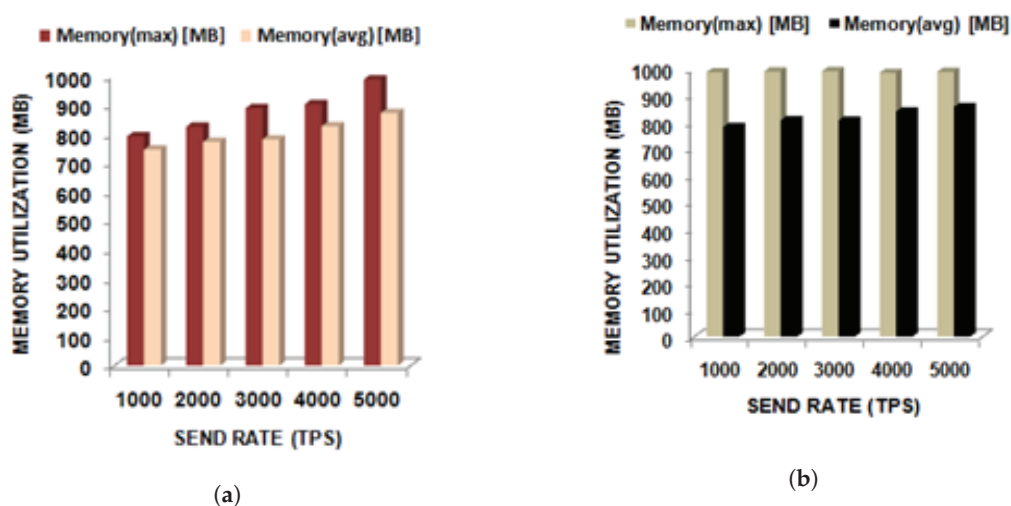


Figure 3. Cont.

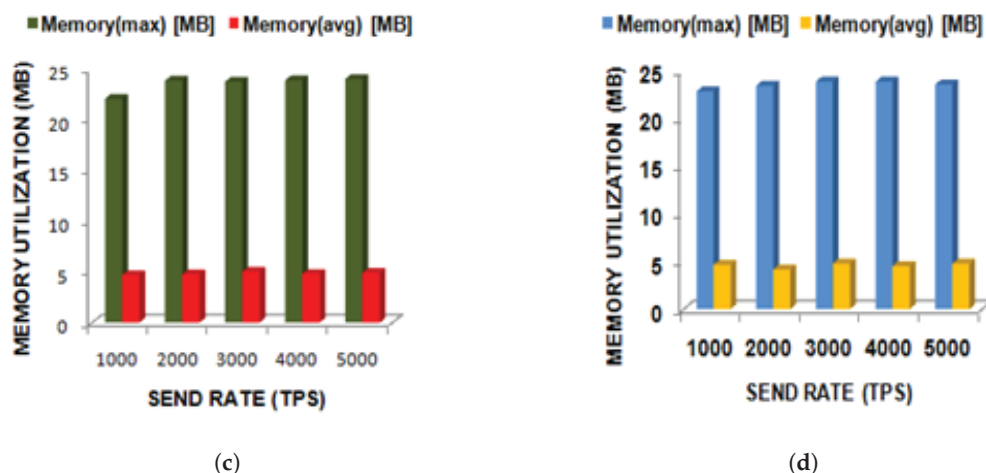


Figure 3. (a) Memory utilization of `createMarketItem` function; (b) memory utilization of `createMarketSale` function; (c) memory utilization of `createToken` function; (d) memory utilization of `fetchItemListed` function. (a) shows the memory utilization (maximum memory used and Average memory utilized) of the smart contract—`createMarketItem` versus send rate (in TPS). (b) shows the memory utilization (maximum memory used and Average memory utilized) of the smart contract—`createMarketSale` versus send rate (in TPS). (c) shows the memory utilization (maximum memory used and Average memory utilized) of the smart contract—`createToken` versus send rate (in TPS). (d) shows the memory utilization (maximum memory used and Average memory utilized) of the smart contract—`fetchItemListed` versus send rate (in TPS).

Figure 4a shows the graph of the CPU utilization in % versus the send rate (1000–5000 TPS). For the `createMarketItem` function, the graph shows that the maximum CPU utilization ranges from 90.90% (4000 TPS) up to 93.95% (1000 TPS). For the `createMarketSale` function, Figure 4b shows the maximum memory utilization is in the range of 26.58% (2000 TPS) up to 29.98% (5000 TPS). For the function `createToken`, Figure 4c shows the maximum memory utilization ranges from 28.0% (1000 TPS) up to 29.8% (5000 TPS). For the `fetchItemListed` function, Figure 4d shows the maximum memory utilization ranges from 23.2% (1000 TPS) up to 28.4% (5000 TPS). In summary, the memory utilization function is highest for `createMarketItem`, within the range of 90.9–93.9%. However, for the other three functions—`createMarketSale`, `createToken`, and `fetchItemListed`—the CPU utilization is low, within the range of 25.7–29.9%.

The memory utilization graph shown in Figure 3 shows that for the `createMarketItem` and `createMarketSale` functions, memory utilization remains high as the send rate increases. Both maximum and average memory utilization are within the range of 800–1000 MB, indicating that this function is resource-intensive and incurs significant processing tasks. The `createToken` and `fetchItemListed` functions have much lower memory usage, with a maximum memory of 25 MB and an average memory of 5 MB, which indicates efficiency in creating tokens with low resource demand, as token creation and tracking waste units is a lighter task. The CPU utilization graph, as shown in Figure 4, depicts that for the `createMarketItem` function, the maximum CPU utilization percentage is high, nearing 90–100%, regardless of the send rate. Average CPU consumption is significantly lower at around 40–50%, which indicates that creating market items for NFTs is CPU-intensive, as it handles complex processes like minting NFTs for waste management and recording metadata. However, the load cannot be sustained continuously. The CPU utilization of the `createMarketSale` and `createToken` functions shows that the minimum CPU consumption is around 25–30%, with an average CPU utilization % below 10%, which suggests that facilitating market sales of NFTs is less CPU-demanding compared to the `createMarketItem` function. In addition, the creation of tokens (representing unique identifiers for waste items) appears to be lightweight, requiring relatively few CPU

resources, while the `fetchItemListed` function also requires minimal CPU resources; this makes the system efficient in token generation and retrieving data, which is critical for large-scale waste management scenarios. The proposed model design is scalable, and this is shown by the fact that, for all the function designs, CPU utilization and memory consumption remain consistent across the send rates, which is vital for blockchain-based waste management systems as the volume of transactions increases. The memory utilization and CPU utilization graphs for the `createMarketItem` function show high resource utilization, as the `MarketItem` is likely a struct with multiple fields and, as we assign a new `MarketItem` to the mapping, it consumes memory (storage) for each of these fields. The `id-market` item is a key value pair in the mapping and requires additional storage. The storage consumption increases significantly as the mapping grows with more entries. Payable addresses, like `payable(msg.sender)` and `payable(address(this))`, are stored in the struct, called `MarketItem`, which adds additional data to the structure, increasing memory usage. For many token listings, if this function is called multiple times, the storage usage increases.

As a result of the above-mentioned processing, it impacts scalability by gas cost, throughput reduction, and network congestion. Hence, off-chain data storage, sharding, and lazy minting will be preferred in our future work.

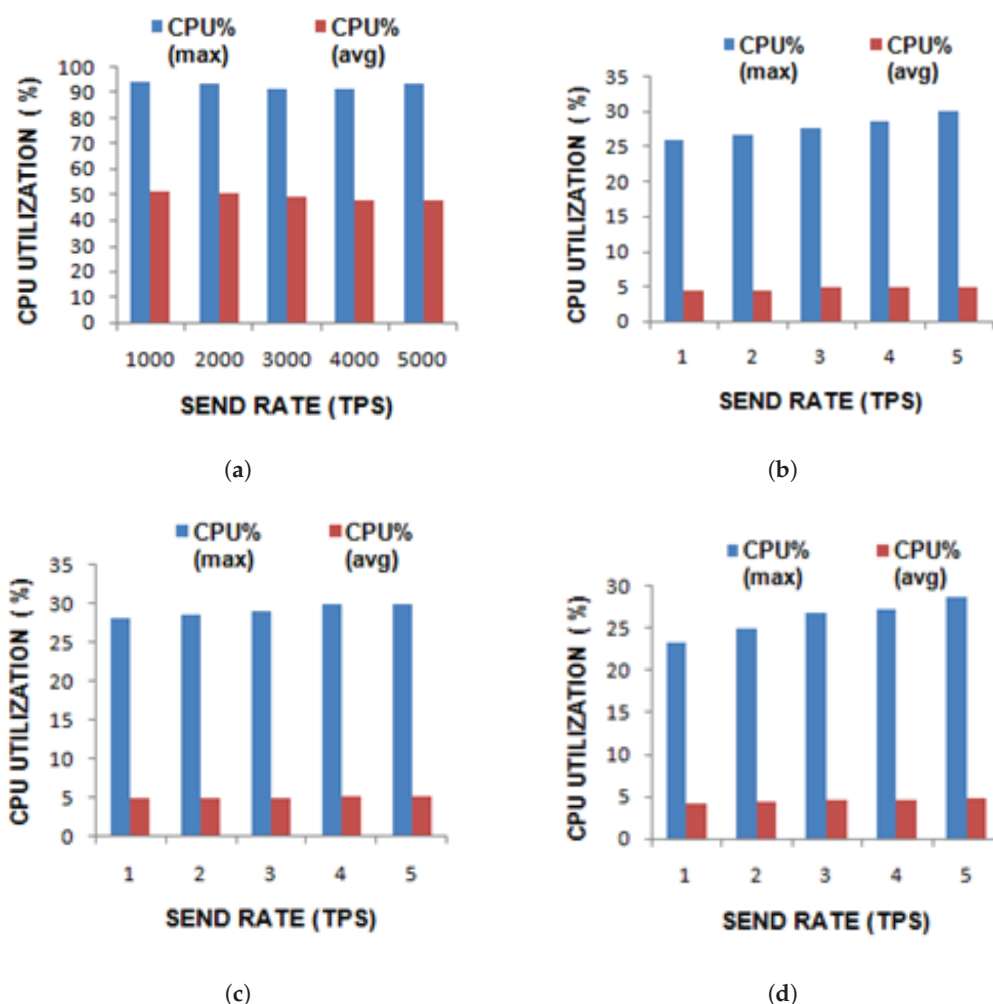


Figure 4. (a) CPU utilization of `createMarketItem` function; (b) CPU utilization of `createMarketSale` function; (c) CPU utilization of `createToken` function; (d) CPU utilization of `fetchItemListed` function. (a) shows the CPU utilization of the smart contract—`createMarketItem` versus send rate (in TPS). (b) shows the CPU utilization of the smart contract—`createMarketSale` versus send rate (in TPS). (c) shows the CPU utilization (maximum memory used and average memory utilized) of the smart contract—`createToken` versus send rate (in TPS). (d) shows the CPU utilization of the smart contract—`fetchItemListed` versus send rate (in TPS).

5. Conclusions

In this article, we have proposed a blockchain-enabled non-fungible token and dynamic bidding using a Walrasian equilibrium to achieve iterative price adjustments for sustainable waste management. We used the ZKP validity mechanism, an innovative framework based on NFTs that facilitates the minting of waste and listing it in the marketplace, as well as the creation of the market item, transferring ownership of funds and items between parties, fetching NFTs, and returning items listed by users. With a seamless interaction, an elegant GUI was developed, with the features of reliability, transparency, and security. Additionally, we conducted a comparative analysis, evaluating performance using the Hyperledger Caliper tool, examining transaction rates, and analyzing the resulting metrics. Analysis of the latency and throughput was performed for the token operations—CreateMarketItem, createMarketSale, createToken, and fetchItemListed. A detailed performance evaluation of our built model was also conducted based on the maximum and minimum CPU utilization, memory utilization, and traffic in/out of each of the token operations (CreateMarketItem, createMarketSale, createToken, fetchItemListed) for each of the send rates from 1000 TPS to 5000 TPS. The results are plotted, proving the efficacy of the current proposed model. The results show that the maximum average latencies of CreateMarketItem, createMarketSale, createToken, and fetchItemListed are 8.76 s, 10.37 s, 11.09 s, and 10.51 s, respectively, and the minimum average latencies are 8.71 s, 9.84 s, 10.89 s, and 10.51 s, respectively. Average latency is relatively high for the CreateToken function, whereas it is lower for the createMarketItem function. Another trend also reveals that as the send rate (transactions per second) increases, the average latency comes down gradually. createMarketItem has the highest throughput, averaging at 251.5 s, and the createToken function has the lowest throughput, averaging at 32.5 s. The average memory consumption of createMarketItem and createMarketSale lies within the range of 772–872 MB for the range of TPS from 1000 to 5000. The average memory consumption of createToken and fetchItemListed lies within the range of 4.6 MB to 5.01 MB for the range of TPS from 1000 to 5000. The traffic in lies within the range of 9.93 to 13.01 and the traffic out ranges from 28.1 MB to 32.6 MB for all the four tokens analyzed. The % maximum CPU utilization of the four functions lies in the range of 23.2 MB to 29.9 MB for the three functions createMarketSale, createToken, and fetchItemListed, and for the createMarketItem, it is 92.95 MB. This is a new dimension in the performance aspect of non-fungible token-enabled sustainable waste management approaches using the zero-knowledge proof (ZKP) validity mechanism. Sustainable waste management is very important, and the advantages of the NFT technology of blockchain, which has also been integrated with the based verification of parties method, have paved the way for a unique dimension in this arena. This promises to bring new developments in waste management and asset ownership through NFTs with the advantages of blockchain.

The limitations of the proposed research would be in providing the participants with real-time data on the environmental impact of their contributions, such as the amount of CO₂ saved or resources recycled. This gives a sense of purpose and also motivates further engagement. Another limitation would be the dynamic adjustment of resource allocation using a load balancer based on real-time network activity, which ensures smooth operation during peak loads. Future work for this research should aim to expand the Greenlink framework's applicability to adjacent domains, like supply chain management, recycling, and carbon trading, to demonstrate its versatility and further validate its utility. Another ongoing work will aim to develop a dynamic load balancer for dynamically adjusting the resource allocation.

Assumption 2. *The allocation of the waste quantity cannot exceed available supply or budgetary limits. Each organization maximizes its utility, while each company maximizes its revenue. Prices are determined based on bidding values and reserve prices.*

Author Contributions: Conceptualization, C.S.K., A.B.P. and A.P.S.; Methodology, C.S.K. and K.H.K.R.; Software, C.S.K. and A.B.P.; Validation, C.S.K.; Formal analysis, A.P.S.; Resources, A.B.P.; Writing—original draft, C.S.K.; Writing—review & editing, K.H.K.R.; Supervision, A.P.S.; Funding acquisition, K.H.K.R. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Data Availability Statement: The original contributions presented in this study are included in the article. Further inquiries can be directed to the corresponding authors.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Khiem, H.G.; Khanh, H.V.; Huong, H.L.; Quy, T.L.; Phuc, T.N.; Ngan, N.T.K.; Triet, M.N.; Bang, L.K.; Trong, D.P.N.; Hieu, M.D.; et al. Implementing a blockchain, smart contract, and NFT framework for waste management systems in emerging economies: An investigation in Vietnam. *Int. J. Adv. Comput. Sci. Appl.* **2023**, *14*, 985–996.
2. Bułkowska, K.; Zielińska, M.; Bułkowski, M. Implementation of blockchain technology in waste management. *Energies* **2023**, *16*, 7742. [CrossRef]
3. Jiang, P.; Zhang, L.; You, S.; Van Fan, Y.; Tan, R.R.; Klemeš, J.J.; You, F. Blockchain technology applications in waste management: Overview, challenges and opportunities. *J. Clean. Prod.* **2023**, *421*, 138466. [CrossRef]
4. Tian, Z. Post-merge carbon footprint analysis and sustainability in the NFT Art Market. *Arts* **2023**, *12*, 211. [CrossRef]
5. Mielcarek, A.; Kłobukowska, K.; Rodziewicz, J.; Janczukowicz, W.; Bryszewski, K., Ł. Water Nutrient Management in Soilless Plant Cultivation versus Sustainability. *Sustainability* **2023**, *16*, 152. [CrossRef]
6. Chiacchio, F.; D’Urso, D.; Oliveri, L.M.; Spitaleri, A.; Spampinato, C.; Giordano, D. A non-fungible token solution for the track and trace of pharmaceutical supply chain. *Appl. Sci.* **2022**, *12*, 4019. [CrossRef]
7. Chiquito-Contreras, R.G.; Hernandez-Adame, L.; Alvarado-Castillo, G.; Martínez-Hernández, M.d.J.; Sánchez-Viveros, G.; Chiquito-Contreras, C.J.; Hernandez-Montiel, L.G. Aquaculture—Production System and Waste Management for Agriculture Fertilization—A Review. *Sustainability* **2022**, *14*, 7257. [CrossRef]
8. El-Rayes, N.; Chang, A.; Shi, J. Plastic Management and Sustainability: A Data-Driven Study. *Sustainability* **2023**, *15*, 7181. [CrossRef]
9. Triet, M.N.; Khanh, H.V.; Huong, H.L.; Khiem, H.G.; Phuc, T.N.; Ngan, N.T.K.; Quy, T.L.; Bang, L.K.; Trong, D.P.N.; Hieu, M.D.; et al. Leveraging Blockchain, Smart Contracts, and NFTs for Streamlining Medical Waste Management: An Examination of the Vietnamese Healthcare Sector. *Int. J. Adv. Comput. Sci. Appl.* **2023**, *14*, 958–970.
10. Ahmad, R.W.; Salah, K.; Jayaraman, R.; Yaqoob, I.; Omar, M. Blockchain for waste management in smart cities: A survey. *IEEE Access* **2021**, *9*, 131520–131541. [CrossRef]
11. Ma, X.; Yuan, H.; Du, W. Blockchain-Enabled Construction and Demolition Waste Management: Advancing Information Management for Enhanced Sustainability and Efficiency. *Sustainability* **2024**, *16*, 721. [CrossRef]
12. Steenmans, K.; Taylor, P.; Steenmans, I. Blockchain technology for governance of plastic waste management: Where are we? *Soc. Sci.* **2021**, *10*, 434. [CrossRef]
13. Al-Farsi, S.; Rathore, M.M.; Bakiras, S. Security of blockchain-based supply chain management systems: Challenges and opportunities. *Appl. Sci.* **2021**, *11*, 5585. [CrossRef]
14. Soori, M.; Jough, F.K.G.; Dastres, R.; Arezoo, B. Blockchains for Industrial Internet of Things in Sustainable Supply Chain Management of Industry 4.0, A Review. *Sustain. Manuf. Serv. Econ.* **2024**, *3*, 100026. [CrossRef]
15. Aslam, J.; Saleem, A.; Khan, N.T.; Kim, Y.B. Factors influencing blockchain adoption in supply chain management practices: A study based on the oil industry. *J. Innov. Knowl.* **2021**, *6*, 124–134. [CrossRef]
16. Dietrich, F.; Ge, Y.; Turgut, A.; Louw, L.; Palm, D. Review and analysis of blockchain projects in supply chain management. *Procedia Comput. Sci.* **2021**, *180*, 724–733. [CrossRef]
17. Duan; Keru; Pang, G.; Lin, Y. Exploring the current status and future opportunities of blockchain technology adoption and application in supply chain management. *J. Digit. Econ.* **2023**, *2*, 244–288. [CrossRef]
18. Dudczyk, P.; Dunston, J.; Crosby, G.V. Blockchain Technology for Global Supply Chain Management: A Survey of Applications, Challenges, Opportunities and Implications. *IEEE Access* **2024**, *12*, 70065–70088. [CrossRef]

19. Marin, O.; Cioara, T.; Todorean, L.; Mitrea, D.; Anghel, I. Review of Blockchain Tokens Creation and Valuation. *Future Internet* **2023**, *15*, 382. [CrossRef]
20. Rodinko, M.; Oliynykov, R.; Nastenka, A. Decentralized Proof-of-Burn auction for secure cryptocurrency upgrade. *Blockchain Res. Appl.* **2024**, *5*, 100170. [CrossRef]
21. Cui, S.; Xu, S.; Hu, F.; Zhao, Y.; Wen, J.; Wang, J. A consortium blockchain-enabled double auction mechanism for peer-to-peer energy trading among prosumers. *Prot. Control Mod. Power Syst.* **2024**, *9*, 82–97. [CrossRef]
22. Kadadha, M.; Mizouni, R.; Singh, S.; Otrok, H.; Ouali, A. ABCrowd an auction mechanism on blockchain for spatial crowdsourcing. *IEEE Access* **2020**, *8*, 12745–12757. [CrossRef]
23. Agarwal, S.; Kapoor, S.; Jain, A. Double Auction Mechanism Based Energy Trading in Blockchain Enabled Smart Microgrid. In Proceedings of the 2024 IEEE 3rd International Conference on Electrical Power and Energy Systems (ICEPES), Faridabad, India, 21–22 June 2024; pp. 1–7.
24. Liu, T.; Wu, J.; Chen, L.; Wu, Y.; Li, Y. Smart contract-based long-term auction for mobile blockchain computation offloading. *IEEE Access* **2020**, *8*, 36029–36042. [CrossRef]
25. Yang, X.; Li, W. A zero-knowledge-proof-based digital identity management scheme in blockchain. *Comput. Secur.* **2020**, *99*, 102050. [CrossRef]
26. Wali, H.; Kulkarni, V.; Ganiger, R.; Iyer, N.C. ZKAV: Zero Knowledge proof for AV. *Procedia Comput. Sci.* **2024**, *237*, 891–898. [CrossRef]
27. Soewito, B.; Marcellinus, Y. IoT security system with modified Zero Knowledge Proof algorithm for authentication. *Egypt. Inform. J.* **2021**, *22*, 269–276. [CrossRef]
28. Zhou, L.; Diro, A.; Saini, A.; Kaisar, S.; Hiep, P.C. Leveraging zero knowledge proofs for blockchain-based identity sharing: A survey of advancements, challenges and opportunities. *J. Inf. Secur. Appl.* **2024**, *80*, 103678. [CrossRef]
29. Kuznetsov, O.; Rusnak, A.; Yezhov, A.; Kanonik, D.; Kuznetsova, K.; Karashchuk, S. Enhanced Security and Efficiency in Blockchain with Aggregated Zero-Knowledge Proof Mechanisms. *IEEE Access* **2024**, *12*, 49228–49248. [CrossRef]
30. Xu, X. Zero-knowledge proofs in education: A pathway to disability inclusion and equitable learning opportunities. *Smart Learn. Environ.* **2024**, *11*, 7. [CrossRef]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

MDPI AG
Grosspeteranlage 5
4052 Basel
Switzerland
Tel.: +41 61 683 77 34

Mathematics Editorial Office
E-mail: mathematics@mdpi.com
www.mdpi.com/journal/mathematics



Disclaimer/Publisher's Note: The title and front matter of this reprint are at the discretion of the Guest Editor. The publisher is not responsible for their content or any associated concerns. The statements, opinions and data contained in all individual articles are solely those of the individual Editor and contributors and not of MDPI. MDPI disclaims responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.



Academic Open
Access Publishing

mdpi.com

ISBN 978-3-7258-4658-0