

Power Choquet-Inspired Integral: A Family of Tunable Aggregation Functions with Application to Fuzzy Rule-Based Classification Systems

G. Lucca, B.L. Dalmazo, R.A. Berri, G.P. Dimuro

C3, Universidade Federal do Rio Grande, Brazil

{giancarlo.lucca, gracializdimuro, rafaelberri, dalmazo}@furg.br

T. C. Asmus

IMEF, Universidade Federal do Rio Grande, Brazil

tiagoasmus@furg.br

M. S. Amorim

PPGMC, Universidade Federal do Rio Grande, Brazil

miqueias.amorim@furg.br

C. Marco-Detchart, H. Bustince

DEIM, Universidad Pública de Navarra, Spain

{cedric.marco, bustince}@unavarra.es

Abstract—The Choquet integral is a well-established aggregation operator capable of modeling interactions among criteria, and it has been successfully employed in fuzzy rule-based classification systems through the use of adaptive fuzzy measures such as the Power Measure. More recently, Choquet-inspired aggregation functions have been proposed as an alternative formulation that reduces the amount of required information while preserving the structural principles of the Choquet integral. In this paper, we introduce the Power Choquet-Inspired Integral (PCII), a novel family of aggregation operators that combines the adaptability of the Power Measure with the Choquet-inspired framework. The proposed approach led to a new aggregation mechanism with reduced informational requirements and improved computational efficiency. An extensive experimental study on 33 benchmark datasets demonstrates that PCII consistently outperforms the classical Power Measure in terms of classification accuracy, while also achieving lower execution times. Statistical analyses confirm the significance of the observed improvements, supporting the effectiveness of the proposed operator in classification tasks.

Index Terms—aggregation function, Choquet integral, Choquet inspired aggregation function, Power measure

I. INTRODUCTION

The discrete Choquet integral [1] is a flexible aggregation operator [2] defined with respect to a fuzzy measure [3]. Its main advantage lies in the ability to explicitly model interactions among criteria, which makes it particularly suitable for multi-criteria decision making and information fusion problems. In contrast to additive aggregation models, the Choquet integral does not assume independence among attributes, allowing for the representation of redundancy and synergy effects between criteria.

Within the context of classification problems [4], the Fuzzy Reasoning Method (FRM) [5] constitutes the core of Fuzzy Rule-Based Classification Systems (FRBCSs) [6]. In [7], a new FRM was introduced in which the Choquet integral is employed to aggregate the information associated with each

class. This formulation enabled the construction of classifiers capable of capturing interactions among features, leading to improvements in classification performance.

In the same work, the authors proposed the Exponential Cardinality, also referred to as the Power Measure (PM). This fuzzy measure is parameterized by an exponent q , which is learned and adjusted for each class of the problem, by a genetical algorithm, providing additional flexibility to the aggregation process.

Motivated by this framework, several generalizations of the Choquet integral have been proposed in the literature. A comparative analysis presented in [8] showed that, among the studied, the PM achieved superior performance in classification tasks, reinforcing its relevance in this domain.

More recently, Choquet-inspired aggregation functions were introduced in [9]. This approach replaces the explicit use of fuzzy measures with general increasing functions that depend directly on the considered inputs, reducing the amount of required information while preserving the fundamental aggregation structure.

Having this in mind, the objective of this paper is to introduce the Power Choquet-Inspired Integrals (PCII), a new family of aggregation operators that combines the Choquet-inspired formulation with the approach of the Power Measure. The proposed operator retains the adaptability of the PM while benefiting from the reduced informational requirements of the Choquet-inspired framework. As a result, PCII provides a class-dependent aggregation mechanism that is both expressive and computationally efficient.

The remainder of this paper is organized as follows. Section II presents the theoretical background. Section III discusses the proposed aggregation operator and its properties. The experimental methodology is described in Section IV. The obtained results are reported in Section V, and the conclusions are drawn in the final section.

This work was partially funded by the following funding agencies: FAPERGS (24/2551-0001396-2, 24/2551-0000723-7, 23/2551-0001865-9), CNPq (304118/2023-0, 407206/2023-0), PID2022-136627NB-I00, MCIN/AEI/10.13039/501100011033/FEDER,UE.

II. THEORETICAL FRAMEWORK

This section recalls the main theoretical concepts related to this study. In what follows, consider $N = \{1, \dots, n\}$, for an arbitrary $n > 0$.

A function $A : [0, 1]^n \rightarrow [0, 1]$ is said to be an aggregation function [10] if **(A1)** it is increasing in each argument and **(A2)** it satisfies the boundary conditions: $A(0, \dots, 0) = 0$ and $A(1, \dots, 1) = 1$.

A t-norm is an aggregation function $T : [0, 1]^2 \rightarrow [0, 1]$ that is associative, commutative and has 1 as neutral element. A t-conorm is an aggregation function $S : [0, 1]^2 \rightarrow [0, 1]$ that is associative, commutative and has 0 as neutral element.

An n -dimensional overlap function is an aggregation function $O_n : [0, 1]^n \rightarrow [0, 1]$ that is symmetric, continuous and respect the following conditions: **(O_n1)** $O_n(x_1, \dots, x_n) = 0 \Leftrightarrow$ there exists $i \in \{1, \dots, n\}$ such that $x_i = 0$; **(O_n2)** $O_n(x_1, \dots, x_n) = 1 \Leftrightarrow x_i = 1$ for all $i \in \{1, \dots, n\}$. An n -dimensional grouping function is an aggregation function $G_n : [0, 1]^n \rightarrow [0, 1]$ that is symmetric, continuous and respect the following conditions: **(G_n1)** $G_n(x_1, \dots, x_n) = 0 \Leftrightarrow x_i = 0$ for all $i \in \{1, \dots, n\}$; **(G_n2)** $G_n(x_1, \dots, x_n) = 1 \Leftrightarrow$ there exists $i \in \{1, \dots, n\}$ such that $x_i = 1$.

A function $m : 2^N \rightarrow [0, 1]$ is said to be a fuzzy measure [1] if (m1) it is increasing: for all $X, Y \subseteq N$, if $X \subseteq Y$, then $m(X) \leq m(Y)$, and (m2) it satisfies the boundary conditions: $m(\emptyset) = 0$ and $m(N) = 1$. Fuzzy measures express the relevance of the coalition in the relationship between the aggregated elements. In this paper, we work with the Power Measure [11], considering its behavior in previous works, as it can be adapted to different classes of classification problems. For $A \subseteq N$, the PM is defined as:

$$m_p(A) = \left(\frac{|A|}{n} \right)^q, \text{ with } q > 0. \quad (1)$$

Observe that the PM can be perceived as a normalized cardinality that is distorted by the action of the automorphism (strictly increasing bijection) $\varphi : [0, 1] \rightarrow [0, 1]$, given by $\varphi(x) = x^q$, with $q > 0$.

Definition 1: [1] Let $m : 2^N \rightarrow [0, 1]$ be a FM. The discrete Choquet Integral (CI) is the function $\mathfrak{C}_m : [0, 1]^n \rightarrow [0, 1]$, defined for all $\mathbf{x} = (x_1, \dots, x_n) \in [0, 1]^n$, by:

$$\mathfrak{C}_m(\mathbf{x}) = \sum_{i=1}^n (x_{(i)} - x_{(i-1)}) \cdot m(A_{(i)}), \quad (2)$$

where $(x_{(1)}, \dots, x_{(n)})$ is an increasing permutation on the input $\mathbf{x}$, that is, $0 \leq x_{(1)} \leq \dots \leq x_{(n)}$, where $x_{(0)} = 0$ and $A_{(i)} = \{(i), \dots, (n)\}$ is the subset of indices corresponding to the $n - i + 1$ largest components of $\mathbf{x}$.

Consider $n \geq 3$. For $i, j \in N$, with $i \neq j$, and $\mathbf{x} \in [0, 1]^n$, denote by $\hat{\mathbf{x}}_{i,j}$ the result of applying the projection $P_{(i,j)} : [0, 1]^n \rightarrow [0, 1]^{n-2}$, omitting components i and j .

Definition 2: [9] Consider $n \geq 3$ and a symmetric and increasing function $F : [0, 1]^{n-2} \rightarrow [0, 1]$. The Choquet-inspired

aggregation function is defined as $\Gamma_F : [0, 1]^n \rightarrow [0, 1]$, given for all $\mathbf{x} \in [0, 1]^n$, by

$$\Gamma_F(\mathbf{x}) = x_{(1)} + \sum_{i=1}^{n-1} (x_{(i+1)} - x_{(i)}) F(\hat{\mathbf{x}}_{(i),(i+1)}), \quad (3)$$

where $(x_{(1)}, \dots, x_{(n)})$ is an increasing permutation on the input $\mathbf{x}$, that is, $0 \leq x_{(1)} \leq \dots \leq x_{(n)}$, where $x_{(0)} = 0$, and $\hat{\mathbf{x}}_{(i),(i+1)}$

III. POWER CHOQUET-INSPIRED INTEGRALS

The main proposal of this paper is to introduce a family of Choquet-Inspired Aggregation Functions, by applying the same automorphism $\varphi(x) = x^q$, with $q > 0$, to the function F in Definition 2. First, let us define this family of functions:

Definition 3: Consider $n \geq 3$ and a symmetric and increasing function $F : [0, 1]^{n-2} \rightarrow [0, 1]$. The Power Choquet-inspired Aggregation function (PCII) is defined as $\Gamma_F^q : [0, 1]^n \rightarrow [0, 1]$, given for all $\mathbf{x} \in [0, 1]^n$ and $q > 0$, by

$$\Gamma_F^q(\mathbf{x}) = x_{(1)} + \sum_{i=1}^{n-1} (x_{(i+1)} - x_{(i)}) (F(\hat{\mathbf{x}}_{(i),(i+1)}))^q, \quad (4)$$

where $(x_{(1)}, \dots, x_{(n)})$ is an increasing permutation on the input $\mathbf{x}$, that is, $0 \leq x_{(1)} \leq \dots \leq x_{(n)}$, where $x_{(0)} = 0$, and $\hat{\mathbf{x}}_{(i),(i+1)}$

Proposition 1: The PCII is a family of Choquet-Inspired Aggregation Functions.

Proof: For any $q > 0$, one can rewrite Equation (4) by considering a function $G : [0, 1]^{n-2} \rightarrow [0, 1]$, such that $G = F^q$, as follows:

$$\Gamma_F^q(\mathbf{x}) = x_{(1)} + \sum_{i=1}^{n-1} (x_{(i+1)} - x_{(i)}) (G(\hat{\mathbf{x}}_{(i),(i+1)})), \quad (5)$$

which is the definition of a Choquet-inspired aggregation function, since G is also a symmetric and increasing function. ■

By Proposition 1, every theoretical finding discussed in [9] is valid for the functions provided by Definition 3, since PCIIs are indeed a family of Choquet-Inspired Aggregation Functions. The main draw of PCIIs resides in the adaptability/tunability that is allowed by the exponent q , in a similar fashion that is observed when defining the Choquet integral coupled with the PM.

Remark 1: $(n - 2)$ -dimensional overlap functions and $(n - 2)$ -dimensional grouping functions are good choices for the function F when defining PCIIs, since the composition of those functions by an automorphism (such as φ) are also $(n - 2)$ -dimensional overlap functions and $(n - 2)$ -dimensional grouping functions, respectively [12]. So, the resulting function F^q preserves all the desirable properties of the chosen function F , with no compromises. This, however, may not hold for t-norms and t-conorms.

In Table I we present the definitions of the functions that were considered in our experiments to act as the function F in Equation (4), each one providing a different PCII.

TABLE I

SOME CHOICES OF FUNCTIONS TO ACT AS F IN A PCII WITH $m = n - 2$

Name (ID)	Definition
Minimum (Min)	$\bigwedge_{i=1}^m x_i$
Product (Prod)	$\prod_{i=1}^m x_i$
Lukasiewicz t-norm (Luk)	$\max(\sum_{i=1}^m x_i - (m - 1), 0)$
Hamacher product (Ham)	$\frac{\prod_{i=1}^m x_i}{\sum_{i=1}^m x_i - \prod_{i=1}^m x_i}$
OB overlap (Ob)	$\sqrt{(\prod_{i=1}^m x_i) \cdot (\bigwedge_{i=1}^m x_i)}$
Geometric mean (GM)	$(\prod_{i=1}^m x_i)^{\frac{1}{m}}$
ODIV overlap (Odiv)	$\frac{\prod_{i=1}^m x_i + \bigwedge_{i=1}^m x_i}{2}$
OmM overlap (OmM)	$(\bigwedge_{i=1}^m x_i) \cdot (\bigvee_{i=1}^m x_i^2)$
Sine overlap (Sin)	$\sin\left(\frac{\pi}{2} \cdot (\prod_{i=1}^m x_i)^{\frac{1}{4}}\right)$
Maximum (Max)	$\bigvee_{i=1}^m x_i^2$
Probabilistic sum (PrSum)	$1 - \prod_{i=1}^m 1 - x_i$
Bounded Lukasiewicz Sum (BL)	$\min(1, \sum_{i=1}^m x_i)$
Einstein Sum (ES)	$\frac{\sum_{i=1}^m x_i}{1 + \sum_{k < l} x_k x_l}$
GB grouping (GB)	$1 - \sqrt{(\prod_{i=1}^m 1 - x_i) \cdot (\bigwedge_{i=1}^m 1 - x_i)}$
Dual of geometric mean (GMD)	$1 - (\prod_{i=1}^m 1 - x_i)^{\frac{1}{m}}$
GDiv grouping (GDiv)	$1 - \frac{\prod_{i=1}^m 1 - x_i + \bigwedge_{i=1}^m 1 - x_i}{2}$
GmM grouping (GmM)	$1 - (\bigwedge_{i=1}^m 1 - x_i) \cdot (\bigvee_{i=1}^m 1 - x_i^2)$
Sine grouping (GSin)	$1 - \sin\left(\frac{\pi}{2} \cdot (\prod_{i=1}^m 1 - x_i)^{\frac{1}{4}}\right)$

IV. EXPERIMENTAL FRAMEWORK

This section describes how the proposed approach is applied and evaluated. For clarity, the experimental framework is organized into three parts. First, we outline the classification problem and the datasets considered. Next, we detail the learning procedure of the proposed method, including the parameter configuration. Finally, we present the statistical tests adopted for performance comparison.

A. Classification process and datasets

A classification problem consists of a set of training examples described by input attributes, which are used to induce a predictive model. The objective is to assign a class label to unseen instances based on the patterns learned during training [13]. As previously stated, in this work the underlying model is a fuzzy classifier, namely FARC-HD.

To assess the proposed approach, a 5-fold cross-validation scheme is employed [14]. Each dataset is randomly divided into five equally sized partitions, each containing approximately 20% of the instances. In each fold, four partitions are used for training and the remaining one for testing. This procedure is repeated five times, so that each partition is used once as the test set. For a fair comparison, we highlight that the partitions are the same ones available at KEEL.

The performance of each method is evaluated on every test partition using the accuracy rate, defined as the ratio between the number of correctly classified instances and the total number of instances. The performance is the averaging accuracy values over the five folds, and this average is taken as the overall result of the algorithm (see Section V).

The evaluation of the PCII is based on the same experimental benchmark adopted in previous generalizations of

TABLE II

SUMMARY OF THE SELECTED DATASETS USED IN THIS ANALYSIS.

Id.	Dataset	#Inst.	#Atts.	#Class	Id.	Dataset	#Inst.	#Atts.	#Class
App	Appendicitis	106	7	2	Pen	Penbased	10,992	16	10
Bal	Balance	625	4	3	Pho	Phoneme	5,404	5	2
Ban	Banana	5300	2	2	Pim	Pima	768	8	2
Bnd	Bands	365	19	2	Rin	Ring	740	20	2
Bup	Bupa	345	6	2	Sah	Saheart	462	9	2
Cle	Cleveland	297	13	5	Sat	Satimage	6,435	36	7
Con	Contraceptive	1473	9	3	Seg	Segment	2,310	19	7
Eco	Ecoli	336	7	8	Shu	Shuttle	58,000	9	7
Gla	Glass	214	9	6	Son	Sonar	208	60	2
Hab	Haberman	306	3	2	Spe	Spectheart	267	44	2
Hay	Hayes-Roth	160	4	3	Tit	Titanic	2,201	3	2
Ion	Ionosphere	351	33	2	Two	Twonorm	740	20	2
Iri	Iris	150	4	3	Veh	Vehicle	846	18	4
Led	led7digit	500	7	10	Win	Wine	178	13	3
Mag	Magic	1,902	10	2	Wis	Wisconsin	683	11	2
New	Newthyroid	215	5	3	Yea	Yeast	1,484	8	10
Pag	Pageblocks	5,472	10	5					

the Choquet integral. A total of 33 datasets from the KEEL repository [15] are considered in the analysis.

Table II summarizes the selected datasets, including their identification (ID), name, number of instances (#Inst.), number of attributes (#Atts.), and number of classes (#Class).

For six datasets (*magic*, *page-blocks*, *penbased*, *ring*, *satimage*, and *twonorm*), a stratified sampling strategy was applied, reducing the training set to 10% of the original instances in order to decrease computational cost. In addition, datasets containing missing values, such as *wisconsin*, were preprocessed by removing incomplete instances to ensure data consistency. This approach follows the same presented in [7].

B. Learning the measure and configuration

The proposed PCII formulation relies on an exponent that is independently adapted for each class of the classification problem. It is important to observe that the learning strategy adopted here is not new. It follows the same procedure introduced in [7] and further described in [16].

The method apply the CHC adaptive search algorithm [17]. The solution encoding is based on a chromosome with one gene per class. All individuals in the initial population are initialized with gene values equal to one. The fitness function is defined by the accuracy rate, and the crossover operator is the BLX scheme centered on the parents' centroid [18]. When population diversity is lost, the algorithm performs a restart while preserving the best individual found so far (elitism).

Regarding parameter configuration, no additional tuning was performed. All parameters were kept at their default values from the original algorithm, available in the KEEL framework¹. Further details on the configuration are provided in [16], [19]–[21]. In addition, we provide a Git repository containing the source code, detailed results, and the original algorithm, where the learning method is described².

C. Statistical tests used in the comparisons

To support the analysis of the experimental results, hypothesis testing procedures recommended in the literature are applied [22], [23]. Since the assumptions required for

¹<https://sci2s.ugr.es/keel/algorithms.php>²<https://github.com/Giancarlo-Lucca/Power-Choquet-Inspired-Integral-A-Family-of-Tunable-Aggregation-Functions-with-Application-to-FRBCS>

parametric tests are not satisfied, non-parametric statistical tests are adopted [24].

First, the aligned Friedman rank test [25] is used to detect statistically significant differences among a group of compared methods. This test assigns ranks to the algorithms across datasets, and the method with the lowest average rank is considered the best-performing one.

In addition, Holm's post-hoc test [26] is applied to verify whether the top-ranked method significantly outperforms the others. This test computes adjusted p-values (APVs) to control the family-wise error rate in multiple comparisons. When the APV is below the predefined significance level α , the null hypothesis of equal performance is rejected.

Finally, pairwise comparisons are conducted using the Wilcoxon signed-rank test [27]. In this case, the null hypothesis assumes that the median of the differences between paired results is zero, indicating the absence of a systematic performance difference between the compared methods.

V. EXPERIMENTAL RESULTS

The results obtained by applying the PCIIs to develop a new FRM to deal with classification problems are described in this section. To facilitate the comprehension of the obtained results, we have split the analysis into three different parts: (i) the analysis of the standard Choquet-inspired functions versus PCIIs, (ii) A comparison of the best PCII versus the normal approach and classical PM and, at the end, (iii) The time consumption analysis.

A. Obtained results by different PCIIs

The results obtained in test dealing with different classification problems are reported in Table III. The rows relate to the considered datasets, the columns represent the Choquet-inspired approaches and each cell contains the mean accuracy among the folds. We have highlighted in **boldface** the largest accuracy obtained among all PCIIs for each dataset. That is, we are analyzing the superior performance for each case.

The first analysis of the obtained results is related with the obtained mean. Precisely, the best performance is related with the PCII based on the *GOMM* a grouping function (79.80%). The second largest is a similar result. Based on the overlap Geometric Mean (79.78%). An interesting result is found with the top three best performance. The approaches Sin and Bounded Lukasiewicz provided the same performance (79.68%), which is an interesting future research problem.

Observing the performance based on families of aggregations, t-norms (Min, Prod, Luk, Ham), Overlap (Ob, GM, Odiv, OmM, Sin), Grouping (GB, GMD, Gdiv, GOMM, GSIN) and t-conorms (Max, PrSum, BL, ES). We have that the minimum t-norm provided the largest performance, however the dual, Maximum, is overperformed by the bounded Lukasiewicz. A similar behavior occurs with the overlap *OmM*. It has a low performance (78.84%), however, the dual, *GOMM* is the largest performance in the study. It seems that functions *F* with disjunctive characteristics tend to perform better than their conjunctive (dual) counterparts.

To statistically validate the observations, a group comparison is performed using the aligned Friedman rank test. The results are presented in Table IV, which reports, for each PCII variant, the corresponding average rank (with lower values indicating better performance) and the adjusted p-value (APV) obtained from Holm's post-hoc procedure.

In this analysis, a significance level of 10% is adopted. When a statistically significant difference is detected between a given method and the control approach, the corresponding APV is underlined in the table.

The *GOMM* variant is selected as the control method, as it exhibits the lowest average rank among all compared approaches. Based on the APV values, the null hypothesis of equivalent performance can be rejected when comparing *GOMM* against five methods, namely OmM, Odiv, GMD, Luk, and GSIN. These results indicate that the *GOMM*-based PCII provides a statistically superior performance with respect to a subset of the evaluated PCII.

B. Comparing the best PCII against the power measure

In this section, we provide a comparison of the best PCII found in the previous analysis against the classical PM. This analysis will state the real impact of the new approach.

In Table V the results are described in a two-block layout. To ease the comprehension of the obtained performance, the results follow the same structure as the one previously used.

As shown in Table V, the average accuracy computed over the 33 datasets favors the proposed PCII approach, which achieves a mean value of 79.80%, compared to 79.20% obtained by PM. A dataset-wise inspection reveals that PM attains higher accuracy in 13 datasets, whereas PCII outperforms PM in 18 datasets. In addition, both approaches yield identical results for two datasets, namely *Bupa* and *Titanic*.

Once again, to avoid conclusions based solely on average performance, a non-parametric pairwise statistical analysis is conducted using the Wilcoxon signed-rank test [27]. A two-tailed test is adopted with a significance level of 0.1. The resulting rank sums are $R^+ = 140.5$ for PM and $R^- = 355.5$ for PCII, with an associated p-value of 0.03. This result indicates a statistically significant difference between the two approaches, supporting the conclusion that the proposed PCII provides a consistent improvement over the classical Power Measure across the evaluated datasets.

C. Execution time analysis

As discussed in Section III, the PCII is defined in terms of the Choquet-Inspired functions [9]. Consequently, the fuzzy measure is defined in terms of a small group of information.

A direct question that is raised is: besides its performance, the time consumption is smaller? The goal of this subsection is to answer it.

The main equipment used in this experiment is powered by an AMD Threadripper PRO 5955WX processor, featuring 16 physical cores and support for 32 simultaneous threads, operating at 4.0 GHz (with boost up to 4.5 GHz) and equipped with 64 MB of cache. This configuration provides strong

TABLE III
ACCURACY MEAN RESULTS OBTAINED IN TEST BY THE DIFFERENT CHOQUET-INSPIRED APPROACHES.

DS	Min	Prod	Luk	Ham	Ob	GM	Odiv	OmM	Sin	GB	GMD	Gdiv	GOMM	GSIN	Max	PrSum	BL	ES
App	87.71	86.80	82.03	83.98	87.71	83.03	83.98	83.94	84.94	80.17	86.80	83.98	83.98	82.08	83.03	83.03	84.94	82.12
Bal	81.44	82.72	77.76	84.32	82.72	88.48	81.92	84.48	84.32	86.40	78.08	85.12	83.84	76.64	85.76	86.08	84.32	86.56
Ban	84.11	85.60	84.74	85.60	86.09	83.68	85.17	84.36	84.25	84.51	85.85	85.49	86.06	83.85	84.89	85.43	84.25	84.75
Bnd	69.35	70.99	67.62	69.11	68.79	69.08	66.05	70.14	72.64	68.84	68.51	68.87	69.03	67.39	69.69	68.26	72.64	68.29
Bup	66.67	62.90	64.93	64.35	66.38	63.48	66.09	65.51	64.06	63.48	65.22	65.80	66.96	60.87	65.22	65.22	64.06	65.51
Cle	57.58	56.21	55.19	54.19	56.89	57.28	53.88	52.85	56.89	54.55	56.90	58.92	56.90	57.21	54.53	53.53	56.89	57.23
Con	51.46	49.22	48.67	51.33	48.54	52.48	51.66	50.04	51.46	52.82	49.42	50.99	51.93	46.23	52.75	53.77	51.46	53.50
Eco	76.51	76.81	77.10	73.82	77.41	79.48	75.91	77.39	78.28	75.02	74.42	77.09	78.58	73.23	76.80	79.18	78.28	78.28
Gla	62.62	63.11	64.52	67.28	60.76	65.90	62.61	60.30	64.99	65.43	64.49	63.08	67.31	63.55	64.50	64.97	64.99	62.17
Hab	72.86	72.21	72.54	74.16	72.86	70.89	74.17	74.49	71.88	72.21	73.83	69.92	72.21	72.22	71.88	72.21	71.88	69.93
Hay	78.72	79.49	78.72	78.72	78.72	81.00	78.72	78.72	78.72	78.72	79.49	78.72	78.72	79.54	78.72	78.72	78.72	78.72
Ion	87.75	87.21	89.48	86.64	87.47	90.04	87.77	88.33	90.33	88.90	87.77	88.04	89.75	88.61	89.19	88.04	90.33	88.04
Iri	92.00	93.33	92.67	92.67	92.67	94.00	93.33	93.33	93.33	94.00	92.00	94.00	93.33	95.33	93.33	94.67	93.33	94.00
Led	68.40	68.60	68.40	68.40	69.60	68.60	68.80	68.20	69.40	68.60	68.40	68.40	69.00	69.80	68.00	69.20	69.40	69.20
Mag	80.18	80.18	77.55	80.02	79.76	79.65	80.02	79.60	80.12	79.76	78.23	79.86	79.49	77.50	79.91	80.13	80.12	79.76
New	92.56	93.49	93.02	93.02	93.02	92.56	94.42	93.95	95.35	93.02	93.49	93.49	93.95	93.02	92.56	95.35	95.35	93.95
Pag	93.80	93.61	94.34	93.61	93.79	94.15	94.34	94.34	93.97	94.16	93.79	94.34	93.98	94.52	94.34	94.16	93.97	94.71
Pen	90.00	90.18	88.55	90.09	88.64	91.45	88.91	88.73	91.00	89.73	89.45	89.55	90.73	88.00	91.00	90.36	91.00	89.45
Pho	82.24	80.55	80.01	83.09	81.29	81.05	82.12	82.77	82.36	82.38	80.13	82.66	81.98	80.16	82.88	82.62	82.36	82.57
Pim	72.53	75.00	73.83	75.38	73.69	76.43	72.52	74.22	74.47	75.13	72.00	74.61	74.74	73.82	75.65	74.21	74.47	74.47
Rin	90.00	90.68	84.32	88.65	88.65	89.46	89.05	90.27	90.14	89.59	83.92	90.95	90.27	79.86	90.27	89.73	90.14	89.59
Sah	71.21	70.12	68.61	67.09	68.82	71.64	68.62	68.61	70.77	72.08	71.22	69.92	72.30	67.52	69.04	68.83	70.77	71.19
Sat	78.85	79.63	78.07	79.32	79.47	79.31	78.53	79.63	80.10	80.09	78.54	79.78	79.62	78.53	79.78	77.45	80.10	79.47
Seg	91.39	90.61	90.48	92.08	91.30	92.34	92.21	89.78	92.60	92.90	90.43	92.34	91.69	90.13	91.43	92.34	92.60	92.25
Shu	96.92	97.52	97.33	97.38	97.75	96.60	97.93	97.84	96.87	98.02	97.61	97.66	98.25	97.66	97.61	97.79	96.87	97.93
Son	76.95	78.85	73.10	78.39	80.33	78.89	75.49	75.03	78.39	79.31	77.44	76.97	76.45	75.97	77.90	77.87	78.39	76.02
Spe	77.87	76.76	78.27	78.64	79.77	77.87	76.02	77.87	77.51	79.73	76.02	77.50	79.37	74.90	76.74	77.88	77.51	75.97
Tit	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87	78.87
Two	86.35	83.78	79.86	84.19	83.65	88.65	83.24	83.38	85.00	87.16	80.81	85.41	86.62	77.57	84.86	85.00	85.00	85.81
Veh	68.20	67.02	62.88	67.26	66.91	67.61	68.80	68.21	67.50	68.21	65.95	66.67	68.68	61.94	67.26	67.97	67.50	66.20
Win	96.06	95.51	92.11	94.33	96.06	94.43	95.48	96.03	94.37	93.25	94.94	97.16	94.90	92.71	94.94	94.95	94.37	94.38
Wis	96.19	95.75	95.32	96.93	96.05	96.34	96.34	95.90	96.34	96.64	95.46	96.49	96.34	94.29	96.49	96.78	96.34	96.93
Yea	56.27	57.14	58.56	57.41	55.73	57.95	57.14	54.72	58.29	56.74	56.27	56.81	57.48	56.27	58.02	56.81	58.29	57.41
Mean	79.20	79.10	77.86	79.10	79.10	79.78	78.79	78.84	79.68	79.41	78.36	79.38	79.80	77.27	79.33	79.44	79.68	79.25

TABLE IV
ALIGN FRIEDMAN RANK TEST WITH HOLM POST HOC TEST AMONG THE PCII.

Approach	Rank	APV	Approach	Rank	APV
GOMM	208.20	-	Ham	298.15	0.29
Sin	213.62	1.00	Min	303.23	0.24
BL	213.62	1.00	Prod	312.97	0.14
GM_1	231.92	1.00	Ob	318.23	0.11
PrSum	251.21	1.00	OmM	324.70	<u>0.07</u>
GB	258.97	1.00	Odiv	341.14	<u>0.02</u>
Max	269.74	0.87	GMD	383.85	<u>0.00</u>
Gdiv	273.03	0.87	Luk	426.42	<u>0.00</u>
ES	278.11	0.78	GSIN	447.89	<u>0.00</u>

performance for multitasking workloads that require intensive processing. With a total of 128 GB of DDR4-3200 ECC RDIMM memory (configured as 2×64 GB). It also integrates an NVIDIA RTX A4500 GPU with 20 GB of ECC GDDR6 memory, which is well suited for machine learning workloads.

In Table VI, the execution time is provided. For each dataset we have the mean time among the five folds and the sum of the times (total time), showing the real time consumption.

The results show that for most datasets, PCII requires less execution time than PM, both in terms of mean and total time. This reduction is consistently observed across small and medium-sized datasets, while for larger datasets the difference becomes more pronounced.

When considering the aggregated results at the bottom of Table VI, PCII achieves a lower overall execution time, with a total of 01:36:57 compared to 01:49:47 for PM. Similarly, the average mean execution time per dataset is reduced from

TABLE V
COMPARISON BETWEEN THE BEST PCII, BASED ON $GOMM$, AND THE CLASSICAL POWER MEASURE.

DS	PM	PCII	DS	PM	PCII
App	80.13	83.98	Pen	90.55	90.73
Bal	82.40	83.84	Pho	82.98	81.98
Ban	86.32	86.06	Pim	73.95	74.74
Bnd	68.56	69.03	Rin	90.95	90.27
Bup	66.96	66.96	Sah	69.69	72.30
Cle	55.58	56.90	Sat	79.47	79.62
Con	51.26	51.93	Seg	93.46	91.69
Eco	76.51	78.58	Shu	97.61	98.25
Gla	64.02	67.31	Son	77.43	76.45
Hab	72.52	72.21	Spe	77.88	79.37
Hay	79.49	78.72	Tit	78.87	78.87
Ion	90.04	89.75	Two	84.46	86.62
Iri	91.33	93.33	Veh	68.44	68.68
Led	68.20	69.00	Win	93.79	94.90
Mag	78.86	79.49	Wis	97.22	96.34
New	94.88	93.95	Yea	55.73	57.48
Pag	94.16	93.98	Mean	79.20	79.80

00:00:40 for PM to 00:00:35 for PCII. These results suggest that the proposed approach not only improves classification performance, as discussed in previous sections, but also offers more efficient computational behavior in practice.

VI. CONCLUSIONS

This paper introduced the Power Choquet-Inspired Integral, a new family of aggregation operators that combines the Choquet-inspired framework with the adaptability of the Power Measure. By replacing the explicit dependence on fuzzy measures with increasing functions while preserving class-dependent learning, the proposed approach reduces the amount

TABLE VI
EXECUTION TIME OF THE PM VS. PCII.S

	PM		PCII	
	Mean Time	Total time	Mean Time	Total Time
App	00:00:06	00:00:29	00:00:05	00:00:27
Bal	00:00:39	00:03:14	00:00:35	00:02:57
Ban	00:01:25	00:07:06	00:01:10	00:05:52
Bnd	00:00:27	00:02:16	00:00:23	00:01:55
Bup	00:00:20	00:01:38	00:00:15	00:01:17
Clv	00:00:26	00:02:11	00:00:21	00:01:47
Cnt	00:01:47	00:08:57	00:01:46	00:08:52
Eco	00:00:21	00:01:47	00:00:17	00:01:25
Gls	00:00:14	00:01:09	00:00:12	00:00:58
Hab	00:00:13	00:01:03	00:00:12	00:01:00
Hay	00:00:05	00:00:27	00:00:05	00:00:27
Ion	00:00:22	00:01:52	00:00:20	00:01:38
Iri	00:00:03	00:00:14	00:00:03	00:00:15
Led	00:00:24	00:01:59	00:00:23	00:01:57
Mag	00:01:20	00:06:38	00:01:16	00:06:19
New	00:00:09	00:00:45	00:00:08	00:00:39
Pag	00:00:24	00:01:59	00:00:19	00:01:33
Pen	00:01:04	00:05:21	00:00:53	00:04:27
Pho	00:01:57	00:09:47	00:01:45	00:08:44
Pim	00:00:49	00:04:07	00:00:46	00:03:50
Rng	00:00:35	00:02:55	00:00:30	00:02:32
Sah	00:00:34	00:02:51	00:00:28	00:02:18
Sat	00:00:45	00:03:43	00:00:40	00:03:18
Seg	00:01:39	00:08:13	00:01:16	00:06:21
Sht	00:00:51	00:04:17	00:00:44	00:03:42
Son	00:00:35	00:02:53	00:00:35	00:02:53
Spc	00:00:28	00:02:22	00:00:26	00:02:09
Tit	00:00:43	00:03:33	00:00:39	00:03:14
Two	00:00:34	00:02:52	00:00:34	00:02:48
Veh	00:00:57	00:04:47	00:00:51	00:04:16
Win	00:00:02	00:00:12	00:00:02	00:00:10
Wis	00:00:29	00:02:27	00:00:25	00:02:06
Yea	00:01:09	00:05:43	00:00:58	00:04:51
Sum	00:21:57	01:49:47	00:19:23	01:36:57
Mean	00:00:40	00:03:20	00:00:35	00:02:56

of required information and maintains flexibility in modeling interactions among attributes.

The evaluation was carried out using a fuzzy classifier applied to 33 benchmark datasets, revealing that PCII systematically outperformed the classical PM in terms of classification accuracy, both on average and in the number of datasets where superior results were achieved. The non-parametric statistical tests confirming that these gains are statistically significant and not due to random effects. A detailed analysis of execution time demonstrated that PCII generally incurs lower computational cost than the PM, a result associated with its reduced informational requirements and simpler aggregation structure.

Future works may include the analysis of distribution, variance, or interpretation of q values across classes/datasets. Additionally, the analysis of different inspired functions is another interesting option, and the study of its behavior in regression and decision-making tasks.

REFERENCES

- [1] G. Choquet, "Theory of capacities," *Annales de l'Institut Fourier*, vol. 5, pp. 131–295, 1953–1954.
- [2] G. Beliakov, A. Pradera, and T. Calvo, *Aggregation Functions: A Guide for Practitioners*. Berlin: Springer, 2007.
- [3] T. Murofushi, M. Sugeno, and M. Machida, "Non-monotonic fuzzy measures and the Choquet integral," *Fuz. Sets Syst.*, vol. 64, no. 1, pp. 73 – 86, 1994.
- [4] R. O. Duda, P. E. Hart, and D. G. Stork, *Pattern Classification (2Nd Edition)*. Wiley-Interscience, 2000.
- [5] O. Cordon, M. J. del Jesus, and F. Herrera, "Analyzing the reasoning mechanisms in fuzzy rule based classification systems," *Mathware and Soft Computing*, vol. 5, no. 2-3, pp. 321 – 332, 1998.
- [6] H. Ishibuchi, T. Nakashima, and M. Nii, *Classification and Modeling with Linguistic Information Granules, Advanced Approaches to Linguistic Data Mining*, ser. Adv. Inf. Proc., Berlin: Springer, 2005.
- [7] E. Barrenechea, H. Bustince, J. Fernandez, D. Paternain, and J. A. Sanz, "Using the Choquet integral in the fuzzy reasoning method of fuzzy rule-based classification systems," *Axioms*, vol. 2, pp. 208–223, 2013.
- [8] G. Lucca, J. A. Sanz, G. P. Dimuro, E. N. Borges, H. Santos, and H. Bustince, "Analyzing the performance of different fuzzy measures with generalizations of the choquet integral in classification problems," in *FUZZ-IEEE*, June 2019, pp. 1–6.
- [9] H. Bustince, R. Mesiar, G. Dimuro, J. Fernandez, J. Lafuente, and B. De Baets, "Choquet-inspired aggregation functions," *Information Fusion*, vol. 126, p. 103499, 2026.
- [10] M. Grabisch, J. Marichal, R. Mesiar, and E. Pap, *Aggregation Functions*. Cambridge: Cambridge University Press, 2009.
- [11] G. Lucca, B. L. Dalmazo, D. G. Ayres, T. Asmus, J. Sartori, H. Santos, E. Borges, B. Bedregal, H. Bustince, and G. P. Dimuro, "dcc-integrals: A generalization of cc-integrals by restricted dissimilarity functions with applications on two different contexts," *Applied Soft Computing*, vol. 181, p. 113517, 2025.
- [12] B. C. Bedregal, G. P. Dimuro, H. Bustince, and E. Barrenechea, "New results on overlap and grouping functions," *Information Sciences*, vol. 249, pp. 148–170, 2013.
- [13] E. Alpaydin, *Introduction to Machine Learning*, The MIT Press, 2010.
- [14] P.-N. Tan, M. Steinbach, and V. Kumar, *Introduction to Data Mining*. Boston: Addison-Wesley Longman Publishing Co., Inc., 2005.
- [15] J. Alcalá-Fdez, L. Sánchez, S. García, M. Jesus, S. Ventura, J. Garrell, J. Otero, C. Romero, J. Bacardit, V. Rivas, J. Fernández, and F. Herrera, "Keel: a software tool to assess evolutionary algorithms for data mining problems," *Soft Computing*, vol. 13, no. 3, pp. 307–318, 2009.
- [16] G. Lucca, J. A. Sanz, G. P. Dimuro, B. Bedregal, M. J. Asiain, M. Elkano, and H. Bustince, "CC-integrals: Choquet-like copula-based aggregation functions and its application in fuzzy rule-based classification systems," *Knowledge-Based Systems*, vol. 119, pp. 32 – 43, 2017.
- [17] L. J. Eshelman, "The CHC adaptive search algorithm: How to have safe search when engaging in nontraditional genetic recombination," in *Foundations of Genetic Algorithms*, G. J. E. Rawlings, Ed. San Francisco: Morgan Kaufmann, 1991, pp. 265–283.
- [18] F. Herrera, M. Lozano, and A. M. Sánchez, "A taxonomy for the crossover operator for real-coded genetic algorithms: An experimental study," *Int. J. Int. Syst.*, vol. 18, no. 3, pp. 309–338, 2003.
- [19] G. Lucca, J. Sanz, G. Pereira Dimuro, B. Bedregal, R. Mesiar, A. Kolesárová, and H. Bustince Sola, "Pre-aggregation functions: construction and an application," *IEEE Transactions on Fuzzy Systems*, vol. 24, no. 2, pp. 260–272, April 2016.
- [20] G. Lucca, J. A. Sanz, G. P. Dimuro, B. Bedregal, H. Bustince, and R. Mesiar, "CF-integrals: A new family of pre-aggregation functions with application to fuzzy rule-based classification systems," *Information Sciences*, vol. 435, pp. 94 – 110, 2018.
- [21] G. Lucca, G. P. Dimuro, J. Fernandez, H. Bustince, B. Bedregal, and J. A. Sanz, "Improving the performance of fuzzy rule-based classification systems based on a nonaveraging generalization of CC-integrals named $C_{F_1 F_2}$ -integrals," *IEEE Trans. Fuzzy Syst.*, vol. 27, no. 1, pp. 124–134, Jan 2019.
- [22] S. García, A. Fernández, J. Luengo, and F. Herrera, "A study of statistical techniques and performance measures for genetics-based machine learning: Accuracy and interpretability," *Soft Computing*, vol. 13, no. 10, pp. 959–977, 2009.
- [23] D. Sheskin, *Handbook of parametric and nonparametric statistical procedures*, 2nd ed. Chapman & Hall/CRC, 2006.
- [24] J. Demšar, "Statistical comparisons of classifiers over multiple data sets," *Journal of Machine Learning Research*, vol. 7, pp. 1–30, 2006.
- [25] J. L. Hodges and E. L. Lehmann, "Ranks methods for combination of independent experiments in analysis of variance," *Annals of Mathematical Statistics*, vol. 33, pp. 482–497, 1962.
- [26] S. Holm, "A simple sequentially rejective multiple test procedure," *Scandinavian Journal of Statistics*, vol. 6, pp. 65–70, 1979.
- [27] F. Wilcoxon, "Individual comparisons by ranking methods," *Biometrics*, vol. 1, pp. 80–83, 1945.