

Development of the fuzzy grid partition methods in generating fuzzy rules for the classification of data set

Murni Marbun, Opim Salim Sitompul, Erna Budhiarti Nababan, Poltak Sihombing

Faculty of Computer Science and Information Technology, Universitas Sumatera Utara, Medan, Indonesia

Article Info

Article history:

Received Nov 30, 2022

Revised Nov 2, 2023

Accepted Dec 7, 2023

Keywords:

Classification accuracy

Fuzzy grid partition

Fuzzy rules

Hybrid method

Rough set

ABSTRACT

The main weakness of complex and sizeable fuzzy rule systems is the complexity of data interpretation in terms of classification. Classification interpretation can be affected by reducing rules and removing important rules for several reasons. Based on the results of experiments using the fuzzy grid partition (FGP) approach for high-dimensional data, the difficulty in generating many fuzzy rules still increases exponentially as the number of characteristics increases. The solution to this problem is a hybrid method that combines the advantages of the rough set method and the FGP method, which is called the fuzzy grid partition rough set (FGPRS) method. In the Irish data, the rough set approach reduces the number of characteristics and objects so that data with excessive values can be minimized, and the fuzzy rules produced using the FGP method are more concise. The number of fuzzy rules produced using the FGPRS method at $K=2$ is 50%; at $K=K+1$, it is reduced by 66.7% and at $K=2$ K, it is reduced by 75%. Based on the findings of the data collection classification test, the FGPRS method has a classification accuracy rate of 83.33%, and all data can be classified.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Opim Salim Sitompul

Faculty of Computer Science and Information Technology, Universitas Sumatera Utara

Padang Bulan 202155 USU, Medan, Indonesia

Email: opim@usu.ac.id

1. INTRODUCTION

The fuzzy rule-based classification system is applicable and appropriate for dealing with classification problems because it can build understandable linguistic models [1]. Fuzzy rule-based classification systems can also build interpretable linguistic models that automatically generate if-then fuzzy rules and can be used to classify new observation [2]. The quality of interpretability is a sought-after trait in different classification tasks. Rule-based systems and fuzzy logic can be employed in the interpretation process for classification purposes. A major limitation of a rule-based system is the potential presence of extensive and intricate rules for categorization, which can sometimes take time to comprehend. Reducing rules is challenging due to several factors [3], [4]. Partition-type fuzzy grid approaches can be used to generate fuzzy rules. Generating fuzzy rules involves partitioning the training data into fuzzy subsets using membership functions and creating fuzzy rules for each grid. Several methods that generate fuzzy rules with the grid partition approach have been made [5]-[10]. The experiment's results indicate a high level of proficiency in classifying and comprehending particular low-dimensional patterns. Nonetheless, when dealing with data that has numerous dimensions, the challenge of increasing the number of rules arises when more characteristics are added. Consequently, there is a substantial increase in the quantity of rules generated.

Integrating fuzzy sets and rough set theory can significantly efficiently decrease the number of rules. This will enable the classification of items into their appropriate categories [11]. To condense attribute sets in extensive intuitive fuzzy information systems, a rule base with minimum size and optimal generation configuration time and storage space can be achieved by utilizing genetic algorithms and intuitive fuzzy rough sets [12]. Several researchers have developed a fuzzy rule generation model with the rough set method as the Kernel intuitionistic fuzzy rough set model (KIFRS) method [13] and proposed the integration of fuzzy logic in the application of the fuzzy rule generation method to the functional resonance analysis (FRAM) to evaluate deicing operations from a systemic plane perspective, where the process of integrating fuzzy logic with a large number of input variables produces many rules. To further refine the proposal, the rough set method is used as a data mining tool to generate and reduce the number of rules in classification [14]-[16].

This study aims to address the sharp increase in the number of fuzzy rules generated in a fuzzy rule-based system, specifically when handling data set classification problems. The study introduces a hybrid approach that combines fuzzy grid partition (FGP) and rough set methods to generate fuzzy rules for the classification of data sets. The grid partition method is derived from the adaptive distributed FGP method, which modifies the rule count according to the desired number of partitions.

2. GENERATING FUZZY RULES

The method proposed as a solution to the problem of generating fuzzy rules in dataset classification is a hybrid method of rough sets and FGPs. In the early stages, the rough set method was used to form fuzzy rules, and then the FGP method was used to generate fuzzy rules.

2.1. Formation of rules with the rough set method

The formation of rules is the initial process in a classification system to deal with problems of imprecision, uncertainty, or incompleteness of information. The formation of rules with the rough set method is to get a short rule estimate from an information table. In rough set terminology, the data table is called the information system. An information system is a data table where attributes label table columns while objects label table rows. Formally, the information system is expressed by a 4-tuple, ie [17]:

$$I = (U, AT, V, f) \quad (1)$$

where U is Universe; AT is a set of attributes; and $\forall a \in AT$, V_a is a domain a and then V is the domain of all attributes

$$V = \bigcup_{a \in AT} V_a \quad (2)$$

f is an information function such that $f(x, a) \in V_a$ for each $x \in U$ and $a \in AT$.

Decision system is data information along with its decision *attributes* d [18],

$$\mathcal{A} = (U, A \cup \{d\}) \quad (3)$$

where $d \notin A$, is a decision attribute that is not considered in A . The elements of $\mathcal{A}$ are called conditional attributes.

In order to obtain a reduced set of necessary characteristics, individuals can evaluate an object's attributes by comparing them with the combined attributes of other objects through the use of the discernibility matrix modulo D . When comparing the attributes of an object and the condition attributes are given priority over the decision attributes. If the attribute values are the same, no result will be produced. However, there will be a result if the compared attribute values are different.

An attribute can be omitted in the rough set without losing its value [19] because redundant attributes will not affect the classification results if removed. Attributes not included in reduce are useless attributes for classifying elements in the universe. In selecting minimal (reduct) attributes from a set of conditional attributes, the prime implicant boolean function is used in the discernibility matrix modulo D . The results are written as the conjunctive normal form (CNF) formula [20]. *Reduct* is a minimal representation of the original data set and is defined as a minimum subset R of set C in such a way as to attribute set D . Suppose the attribute set is called the reduct of C if $T'=(U, R, D)$ is independent and [21],

$$\gamma_R(D) = \gamma_C(D) \quad (4)$$

That R is a minimum subset if,

$$\gamma_{-(R - \{a\})}(D) \neq \gamma_C(D), \forall a \in R \quad (5)$$

No attribute can be removed from the subset without affecting the dependency level. A specific set of data may have multiple sets of reductions, and the collection of all reductions is represented as (6):

$$R_{ALL} = \{ X | X \subseteq C, \gamma_R(D) = \gamma_C(D) \text{ and } \gamma_{R-\{a\}}(D) \neq \gamma_C(D), \forall a \in X \} \quad (6)$$

The classic rough set model is susceptible to the ambiguity of information. This causes the approximation set in the form of β -lower and β -upper approximation not to be appropriately determined. Therefore, Ziarko modified the rough set theoretical by introducing the variable precision rough set (VPRS) method. The VPRS model is a development model of the classic rough set method intended to handle information uncertainty with the additional assumption of the relative degree of error in classification. The relative degree of misclassification is denoted as the precision value β , which is the main part before determining the reduced variable. The relative degree of misclassification is denoted as (7) [22]:

$$c(C, Y) = \begin{cases} 1 - \frac{|C \cap Y|}{|C|}; & |C| > 0 \\ 0 & |C| = 0 \end{cases} \quad (7)$$

C is included in the condition attribute, a component of the decision set Y. If the number of members in both groups is the same, then the error rate is zero. If the error rate exceeds zero, the object is a potential candidate for β -reduct. The value of β -reduct is determined by choosing the condition that results in the lowest relative degree of error, as calculated by (8) and (9):

$$\beta(C, Y) = \max(m_1, m_2) \quad (8)$$

with

$$m_1 = 1 - \min(c(C, Y) > 0.5) \quad (9)$$

$$m_2 = \max(c(C, Y) < 0.5)$$

2.2. Generation of fuzzy rules with grid fuzzy partition

Grid-type fuzzy partitions have been frequently utilized in fuzzy control since 1970. These partitions offer enhanced interpretability for two-dimensional issues, such as two-dimensional design and single input-output systems that rely on fuzzy rules, which holds especially when dealing with large datasets [23]:

- The rigidity of making changes to membership functions. Due to the utilization of multiple fuzzy rules for each antecedent fuzzy set, enhancing the precision of one fuzzy rule by modifying the membership function may reduce the accuracy of several other fuzzy rules.
- The increase in the number of input variables results in an exponential growth in fuzzy rules. Let's assume that L represents the quantity of fuzzy set antecedents corresponding to each variable n. In the problem with n dimensions, the total number of fuzzy grids can be defined as L raised to the power of n.

This difficulty can be eliminated by assigning different fuzzy set antecedents to each fuzzy rule, as shown in Figure 1, where each fuzzy rule has its own antecedent fuzzy set. That is, multiple fuzzy rules share no fuzzy set antecedents. This fuzzy rule-based system with these fuzzy rules is called a descriptive model of grid-type fuzzy rule-based systems [24].

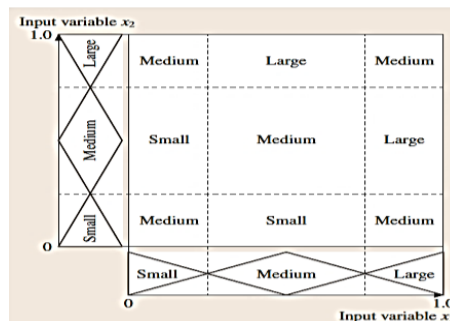


Figure 1. A descriptive model grid type fuzzy rule

Assume that there is m training patterns $X_p = (x_{p1}, \dots, x_{pn})$, $p=1, 2, \dots, m$ of M different classes where X_p is an attribute n-dimensional vector where x_{pi} is the i-th attribute value of the i-th training pattern ($i=1, 2, \dots, n$). In the classification problem for an M-class with n dimensions, the if-then fuzzy rules are as follows [8]:

Rp : If x_1 is A_{p1} and... and x_n is A_{pn} then class C_{qp} with CF_p

where Rp is the label of the p -th fuzzy rule, $X=(x_1, \dots, x_n)$ is the n dimension of the vector of a pattern, A_{pi} is a fuzzy set of antecedents, and C_p is a class label

Suppose that every fuzzy pattern space is subset partitioned to $K \{A_1^K, A_2^K, \dots, A_K^K\}$ where A_i^K described by the triangular membership function, then:

$$\mu_i^K(x) = \max\{1 - |x - a_i^K|/b^K, 0\} \quad (10)$$

where:

$$i = 1, 2, \dots, K$$

$$a_i^K = (i - 1)/(K - 1) \quad (11)$$

$$b^K = \frac{1}{K-1} \quad (12)$$

where K is the number of partitions, a_i^K is the center point for each of the triangle functions, and b^K is the distance from the center point to the legs of the triangle function.

The following steps determine the conclusion C_{ij}^K and degree level CF_{ij}^K of the rule:

a. Count β_{Ct} for $t=1, 2, \dots, M$

$$\beta_{Ct} = \sum_{p \in Ct} \mu_i^K(x_{p1}) \cdot (\mu_j^K(x_{p2})) \quad (13)$$

b. Define class X (CX)

$$\beta_{Cx} = \max \{\beta_{C1}, \beta_{C2}, \dots, \beta_{CM}\} \quad (14)$$

If two or more classes are taken for the maximum value in (14), then the fuzzy rules are in the fuzzy space $A_i^K \times A_j^K$ not generated. Conclusion C_{ij}^K others are defined as CX.

c. If one class gets the maximum value, the rule weight CF_{ij}^K determined in a way:

$$CF_{ij}^K = \frac{|\beta_{Cx} - \beta|}{\sum_{t=1}^M \beta_{Ct}} \quad (15)$$

where:

$$\beta = \sum_{Ct \neq CX} \beta_{Ct} / (M - 1) \quad (16)$$

Conclusion C_{ij}^K is determined from the class that has the largest number $(\mu_i^K(x_{p1}) \cdot (\mu_j^K(x_{p2}))$. The degree degree CF_{ij}^K is in the interval value $[0,1]$. If all patterns $A_i^K \times A_j^K$ come from Class X (CX), then $CF_{ij}^K = 1$. Conversely, if the pattern comes from another class, then $CF_{ij}^K < 1$. In the step, if-then K^2 fuzzy rules are generated from the training pattern in the pattern space $[0,1] \times [0,1]$. The set of generated fuzzy rules is called S_R .

$$S_R = \{R_{ij}^K | i = 1, 2, \dots, K; j = 1, 2, \dots, K\} \quad (17)$$

2.3. Classification with grid partition

The classification process is carried out to determine the accuracy of the rules that have been generated. The steps in carrying out the classification are as (18) and (19) [25]:

a. Calculate the value of α for each class α_{Ct} for $t=1, 2, \dots, M$ using (18), where α is the result of adding the α -predicate to the weight rule

$$\alpha_{Ct} = \max \{(\mu_i^K(x_{p1}) \cdot (\mu_j^K(x_{p2})) \cdot CF_{ij}^K | C_{ij}^K = Ct; R_{ij}^K \in S_R\} \quad (18)$$

b. Determine the maximum value of α (α_{CX}) in each class

$$\alpha_{CX} = \max \{\alpha_{C1}, \alpha_{C2}, \dots, \alpha_{CM}\} \quad (19)$$

If the maximum α_{CX} value is more than one class, the final conclusion cannot be determined, which means that the object cannot be classified (unclassified).

3. PROPOSED METHOD

The development of the fuzzy rule generation method in this research involves four stages in the approach. The initial step consists of preparing the dataset and forming rules using rough set theory. Next, we generate fuzzy rules using the adapted FGP method. Finally, test the method by classifying the data set to assess the accuracy of the classification. The method development architecture proposed in this research can be seen in Figure 2.

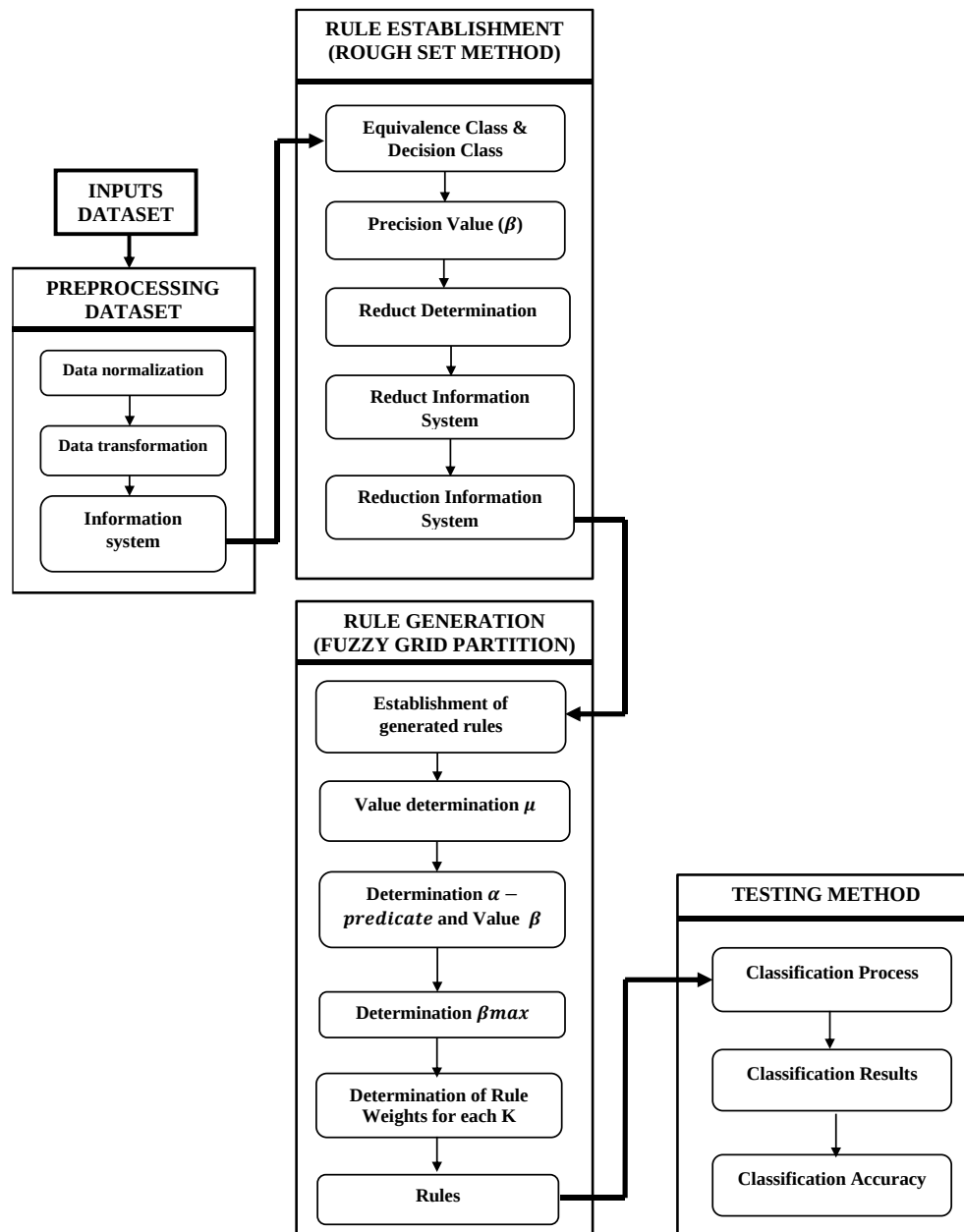


Figure 2. The proposed method architecture

3.1. Input data set

The data in this research comes from the UC Irvine machine learning repository, a dataset the machine learning community uses to assess machine learning algorithms. The training data used is the Iris dataset, which includes four attributes representing conditions: petal length, petal width, sepal length, and sepal width. Apart from that, three attributes represent the decision class, namely setosa, versicolor, and virginica. Normalizing the values in a data set involves transforming the intervals of the actual values in the data set into the range of $[0,1]$. To obtain a new value, it can calculate based on (20):

$$New Value = \frac{Data\ value - smallest\ data\ value}{largest\ data\ value - smallest\ data\ value} \quad (20)$$

As an example of calculating the normalization of the iris data set on data record number one on the sepal length attribute, with the largest data being 7.9 and the smallest data being 4.3. Then the new value, according to (20) is:

$$New Value = \frac{5.1 - 4.3}{7.9 - 4.3} = \frac{0.8}{3.6} = 0.2222$$

Data transformation can be performed on condition attributes and decision attributes. Attributes in this study were only conditional, while decision attributes were not transformed because the data was already in a definite form. The condition attribute is transformed into several set categories. The attribute transformation of the sepal length and petal length conditions is divided into three sets of categories, namely:

- a. Short set category (soca) : ≤ 0.3333
- b. Medium set category (meca) : $0.3333 < meca \leq 0.6667$
- c. Long set category (loca) : > 0.6667

Transformation of sepal width and petal width attributes is divided into three sets of categories, namely:

- a. Small category (smaca) : ≤ 0.3333
- b. Medium category (meca) : $0.3333 < meca \leq 0.6667$
- c. Width category (wica) : > 0.6667

3.2. Formation of rules

Processing the initial data for the formation of rules using the rough set method, the table information system, and the decision system are subject to symbol changes. Data records are converted into objects with variable, 1, 2, 3,..., Xn_n and decisions are changed to letters of the alphabet, where St is for the setosa class, Vc is for the versicolor class, and Vn is for the Virginia class. The following is an excerpt of 9 data records for the Iris dataset from Table 1 with the object names 1, 2, 3, 4, 5, 6, 7, 8, and 9, which are used as sample objects in forming rules. Data representation can be seen in Table 1.

Table 1. Information and decision system

Object	Petal lenght	Petal width	Sepal length	Sepal width	Decision
1	soca	meca	soca	smaca	St
2	soca	meca	soca	smaca	St
3	soca	smaca	soca	smaca	St
4	meca	smaca	meca	meca	Vc
5	soca	smaca	meca	meca	Vc
6	soca	smaca	meca	meca	Vc
7	soca	smaca	meca	meca	Vn
8	meca	meca	loca	wica	Vn
9	meca	meca	loca	wica	Vn

The information in Table 1 is organized into equivalence classes (EC) based on certain conditions, and decision classes (DC) are formed by grouping similar decision attributes. The process of calculating the relative degree value is used in (7). Table 2 shows the degree of relationship for each decision in the equivalence class as follows.

Table 2. The degree of relationship

EC	The degree of relationship		
E1	$c(E1, D1) = 0$	$c(E1, D2) = 1$	$c(E1, D3) = 1$
E2	$c(E2, D1) = 0$	$c(E2, D2) = 1$	$c(E2, D3) = 1$
E3	$c(E3, D1) = 1$	$c(E3, D2) = 0$	$c(E3, D3) = 1$
E4	$c(E4, D1) = 1$	$c(E4, D2) = \frac{1}{3}$	$c(E4, D3) = \frac{1}{3}$
E5	$c(E5, D1) = 1$	$c(E5, D2) = 1$	$c(E5, D3) = 0$

The β -reduction value using (8) and (9) is $1/3$; it is located on the EC row E4, so the β -reduction set ($\underline{C}\beta$) is $E4 = \{5, 6, 7\}$. The set of β -reduct ($\underline{C}\beta$) is $E4 = \{5, 6, 7\}$, then the discernibility matrix modulo D on β -reducts for $5 = c \cup d$, $6 = c \cup d$ and $7 = \cap (c \cup d)$. The results of comparing the attribute values

are then converted into the CNF with a simplification are $R_{c\beta} = (a \cup c) \cap (c \cup d)$ where the attribute set consists of $\{a\}, \{c\}, \{c\}, \{d\}$.

From the reducing set β , the attribute with the repeated reduced frequency of occurrence is attribute c . Thus, the minimal attribute subset with minimal and non-recurring frequency of occurrence consists of features a , b , and d . In the post-reduction information system in Table 1, there are still aspects of object redundancy that need to be reviewed. If we look at the redundancy pattern, it can be said that object 1 with 2, object 5 with 6, and object 8 with 9 are redundant. This way, objects with redundant attributes can be reduced to a single rule to derive a decision. The information system table after reduction can be seen in Table 3.

3.3. Generation of fuzzy rules

The initial process to generate the rules is that the data set in the information system reduction in Table 4 is returned to the initial normalized form of the data. Table 5 is the normalized reduction information system which rules will be generated with five data records and three condition attributes.

Table 3. Information system after reduction

Object	Sepal length	Sepal width	Petal length	Petal width	Decision
1	soca	meca	soca	smaca	St
3	soca	smaca	soca	smaca	St
4	meca	smaca	meca	meca	Vc
5	soca	smaca	meca	meca	Vc
7	soca	smaca	meca	meca	Vn
8	meca	meca	loca	wica	Vn

Table 4. Normalized reduction information system

Object	Sepal length	Sepal width	Petal width	Decision
1	0.029	0.375	0.029	St
3	0.057	0.125	0.057	St
4	0.557	0.208	0.557	Vc
5	0.418	0.250	0.418	Vc
7	0.168	0.208	0.168	Vn
8	0.418	0.333	0.418	Vn

Table 5. Rule weight (CF_{ijk}^K) on $K \leftarrow 2$

R_{ijk}^K	Decision	CF_{ijk}^K
R_{111}^2	St	0.437
R_{112}^2	Vn	0.424
R_{121}^2	St	0.578
R_{122}^2	Vn	0.542
R_{211}^2	Vc	0.511
R_{212}^2	Vn	0.258
R_{221}^2	Vc	0.587
R_{222}^2	Vc	0.492

By entering all linguistic variables at partition $K=2$, eight rules can produce a dataset in Table 5, with some attributes of 3. We get $K^d=2^3=8$, where d is the number of dataset attributes. For $K=3$, we get $K^d=3^3=27$, and the candidate rule for $K \leftarrow 4$ is $4^3=64$. The calculation of μ at K for the dataset for each attribute of the object will be determined by (21):

$$\mu_i^K(x) = \frac{b-x}{b-a} \quad (21)$$

Where $\mu_i^K(x)$ is x membership degree; x is object value for each attribute; b is upper limit; and a is lower limit.

The calculation of μ at $K \leftarrow 2$ for the data set with the sepal length attribute with object 1 is as follows:

$$\mu_1^2(0.0278) = \frac{1.000-0.028}{1.000-0.000} = 0.972$$

$$\mu_2^2(0.0278) = \frac{0.028-0.000}{1.000-0.000} = 0.028$$

The next step is to determine the rule weight of CF_{ijk}^K using (13)-(16). The rule weight of CF_{ijk}^K on $K \leftarrow 2$ can be seen in Table 5. In the same way, the value of μ , at $K \leftarrow 2$, μ at $K \leftarrow K+1$, μ at $K \leftarrow 2K$ for each attribute is obtained and the rule weight of CF_{ijk}^K .

The formation of the grid structure using the adapted grid partition technique is carried out based on rule weights. Rule generation is carried out in 2 stages. In the first stage, a grid will be formed as smoothly as possible, provided that all possible rules can be generated, or the iteration will stop when there is a rule that weights one. Meanwhile, in the second stage, new rules will be generated with higher weights to increase the accuracy of these rules [6]. The steps for the adapted grid partition method to obtain the rules are obtained:

- For phase $\leftarrow 1$, the set of S_R rules is obtained, $S_R = \{R_{111}^2, R_{112}^2, R_{121}^2, R_{122}^2, R_{211}^2, R_{212}^2, R_{221}^2, R_{222}^2\}$.

- b. From the set of rules obtained in step 1, it is obtained that the CF value does not have a value of 1. Thus, the next step is the process $K \leftarrow K+1$ or $K \leftarrow 3$.
- c. For $K \leftarrow 3$, the stage has not been added, so the contents of the original S_R are deleted, then replaced with a new one, so the rules generated in this state (Stage $\leftarrow 1$, $K \leftarrow 3$) are $S_R = \{R_{111}^3, R_{112}^3, R_{113}^3, R_{121}^3, R_{122}^3, R_{123}^3, R_{211}^3, R_{212}^3, R_{213}^3, R_{221}^3, R_{222}^3, R_{223}^3, R_{312}^3, R_{313}^3, R_{322}^3, R_{323}^3\}$.
- d. The S_R set obtained does not have all the rules potentially generated because there is a CF=0 value and there is a β_{max} value of 0 in the rules $R_{131}^3, R_{132}^3, R_{133}^3, R_{231}^3, R_{232}^3, R_{233}^3, R_{311}^3, R_{321}^3, R_{331}^3, R_{332}^3$ and R_{333}^3 so that the S_R set is not fulfilled. S_R is not fulfilled then $K \leftarrow 3$ becomes $K \leftarrow 2$ then the set of S_R returns $S_R = \{R_{111}^2, R_{112}^2, R_{121}^2, R_{122}^2, R_{211}^2, R_{212}^2, R_{221}^2, R_{222}^2\}$.
- e. The step is continued to the second stage by paying attention to the set of S_R where the rule a is chosen because it has The minimum rule weight in the S_R . Thus, K and K_{phase1} values are initialized from R_{212}^2 , and $K \leftarrow 2K$ becomes $K=4$. Thus, the rules that have the potential to be generated from $K \leftarrow 2K$ are $R_{313}^4, R_{314}^4, R_{323}^4, R_{324}^4, R_{413}^4, R_{414}^4, R_{423}^4, R_{424}^4$. The results of the recapitulation of the rules for generating fuzzy rules are shown in Table 6.

Table 6. Recapitulation of generating fuzzy rules R_{212}^2 on $K \leftarrow 2K$

R_{ijk}^K	$\Sigma \{\mu_i^K (sl), \mu_j^K (sw), \mu_k^K (pw)\}$			β_{max}	Decision	CF_{ijk}^K
	β_{setosa}	$\beta_{versicolor}$	$\beta_{virginica}$			
R_{313}^4	0.000	0.188	0.000	0.188	Vc	1
R_{314}^4	0.000	0.000	0.000	0.000	-	0
R_{323}^4	0.000	0.313	0.031	0.313	Vc	0.86
R_{324}^4	0.000	0.000	0.219	0.219	Vn	1
R_{413}^4	0.000	0.000	0.000	0.000	-	0
R_{414}^4	0.000	0.000	0.000	0.000	-	0
R_{423}^4	0.000	0.000	0.000	0.000	-	0
R_{424}^4	0.000	0.000	0.000	0.000	-	0

- f. The determination of the rules generated according to Table 6 from R_{212}^2 on $K \leftarrow 2K$ are rules that do not have a value CF=0, namely $S_R = \{R_{313}^4, R_{323}^4, R_{324}^4\}$.
- g. Thus, the rule generated at Stage $\leftarrow 2$, where $K=2K$ is $S_R = \{R_{111}^2, R_{112}^2, R_{121}^2, R_{122}^2, R_{211}^2, R_{212}^2, R_{221}^2, R_{222}^2, R_{313}^4, R_{323}^4, R_{324}^4\}$ where S_R is the set of rules.

3.4. Testing

The process of testing a method involves classifying objects in a data set based on the rules obtained. The classification process can be carried out by (18) and (19).

- a. Calculates the value of α for rule R_{111}^2 for setosa class with object 1.

$$\begin{aligned}\alpha &= \mu_1^2 (sl) \cdot \mu_1^2 (sw) \cdot \mu_1^2 (pw) \cdot CF_{111}^2 \\ &= (0.972) (0.625) (0.958) (0.454) \\ &= 0.264\end{aligned}$$

- b. Determine the largest α value in each class α_{setosa} , $\alpha_{versicolor}$, and $\alpha_{virginica}$. The calculation results for the object 1 in the setosa class, versicolor class, and virginica class can be seen in Tables 7-9.

Table 7. The maximum value of α_{setosa}

R_{ijk}^K	$\mu_i^K (sl)$	$\mu_j^K (sw)$	$\mu_k^K (pw)$	CF_{ijk}^K	α	α_{max}
R_{111}^2	0.972	0.625	0.958	0.454	0.264	0.264
R_{121}^2	0.972	0.375	0.958	0.578	0.202	

- c. The α_{max} value in object 1 is as follows

$$\begin{aligned}\alpha_{max} &= \max \{ \alpha_{setosa}, \alpha_{versicolor}, \alpha_{virginica} \} \\ \alpha_{max} &= \max \{ 0.264 ; 0.009 ; 0.011 \} \\ \alpha_{max} &= 0.264 \text{ (Setosa class)}\end{aligned}$$

Because the maximum α value is in the setosa class, the classification results for object 1 are correct.

Table 8. The maximum value of $\alpha_{versicolor}$

R_{ijk}^K	μ_i^K (sl)	μ_j^K (sw)	μ_k^K (pw)	CF_{ijk}^K	α	α_{max}
R_{211}^2	0.028	0.625	0.958	0.511	0.009	0.009
R_{221}^2	0.028	0.375	0.958	0.585	0.006	
R_{313}^4	0.000	0.000	0.000	1.000	0.000	
R_{323}^4	0.000	0.8749	0.000	0.863	0.000	

Table 9. The maximum value of $\alpha_{virginica}$

R_{ijk}^K	μ_i^K (sl)	μ_j^K (sw)	μ_k^K (pw)	CF_{ijk}^K	α	α_{max}
R_{112}^2	0.972	0.625	0.042	0.424	0.011	0.011
R_{122}^2	0.972	0.375	0.042	0.542	0.008	
R_{212}^2	0.028	0.625	0.042	0.258	0.000	
R_{222}^2	0.028	0.375	0.042	0.492	0.000	
R_{324}^4	0.000	0.875	0.000	1.000	0.000	

4. RESULT AND DISCUSSION

The formation of rules is the stage of processing datasets that have been normalized and transformed before the process of generating rules is carried out. The formation of the rules applies the rough set method to produce a reduction in the attribute set by considering the variable precision or error rate to obtain a reduction in the attribute. Information system reduct in Table 1 is reviewed again about the redundancy pattern of condition attribute values and the resulting target attribute values. In Table 1, objects 1 and 2 are repeating objects, so one can be omitted, with objects 5 with 6 and 8 with 9. The results of forming rules with attributes with minimum reduced frequencies and reduction results can be seen in Table 4.

By including all linguistic variables at $K=2$, 8 rules have the potential to be generated in the dataset in Table 4, with a total of 3 attributes, $K^d=2^3=8$, where d is the number of dataset attributes. $K \leftarrow K+1$ with three attributes, $K^d=3^3=27$ rules, and $K \leftarrow 2K$ with three attributes, $K^d=4^3=64$ rules. The development of the FGP method is the fuzzy grid partition rough set (FGPRS) method to generate fuzzy rules in classifying data sets, reduce the fuzzy rules by 50% for $K=2$, 66.7% for $K \leftarrow K+1$, and 75% for $K \leftarrow 2K$. Table 10 shows the percentage of reduced fuzzy rules.

Table 10. The percentage reduction of the number of fuzzy rules generated

K	Number of generated fuzzy rules		Percentage reduction in the number of fuzzy rules
	FGP	FGPRS	
2	16	8	50
$K+1$	81	27	66.7
$2K$	256	62	75

The FGPRS method produces 11 rules to classify nine objects with four attributes. Based on the results of classification testing, the rules produced using the hybrid method show correct and incorrect classification results. Table 11 shows that 5 data can be classified correctly, and one is classified incorrectly. The classification accuracy percentage is calculated by comparing the correct classifications with the overall classification. Testing with the FGPRS method produces a classification accuracy rate of 83.33%.

Table 11. Classification accuracy

Object	Maximum				Decision	Results
	α_{setosa}	$\alpha_{versicolor}$	$\alpha_{virginica}$	α_{max}		
1	0.264	0.009	0.011	0.264	<i>St</i>	Correctly classified
3	0.344	0.023	0.029	0.344	<i>St</i>	Correctly classified
4	0.067	0.270	0.087	0.270	<i>Vc</i>	Correctly classified
5	0.108	0.067	0.085	0.108	<i>St</i>	Incorrect classified
7	0.010	0.023	0.187	0.187	<i>Vn</i>	Correctly classified
8	0.007	0.006	0.219	0.219	<i>Vn</i>	Correctly classified

5. CONCLUSION

This research produces a method for generating fuzzy rules in classifying data sets. Applying the rough set method in forming rules at the beginning can reduce the number of excessive attributes and similar object data so that the rule formation process becomes more concise. Next, the fuzzy grid partition method produces fuzzy rules, where the number of partitions will influence the number of rules produced. Applying

an adapted grid partition based on the minimum rule weight produces fuzzy rules that do not increase exponentially. The FGPRS method reduces the resulting fuzzy rules by 50% for $K=2$, 66.7% for $K \leftarrow K+1$, and 75% for $K \leftarrow 2K$. Testing the method on iris data resulted in a classification accuracy rate of 83.33%, and all data could be classified. Some suggestions that can be made for further research development in generating fuzzy rules on data sets are to apply other methods at an early stage in forming rules that have the potential to be generated. Other approaches related to forming a grid structure that will influence the number of rules generated also need to be considered

ACKNOWLEDGEMENTS

The author thanks the Kementerian Pendidikan, Kebudayaan, Riset, and the Teknologi Republic of Indonesia for the Doctoral Dissertation Research grant for the 2022 fiscal year with contract number 76/UN5.2.3.1/PPM/KP-DRTPM/L/2022.

REFERENCES

- [1] Y. X. Zhang, X. Y. Qian, J. Wang, and M. Gendeel, "Fuzzy rule-based classification system using multi-population quantum evolutionary algorithm with contradictory rule reconstruction," *Applied Intelligence*, vol. 49, no. 11, pp. 4007–4021, 2019, doi: 10.1007/s10489-019-01478-5.
- [2] A. Borgi, R. Kalai, and H. Zgaya, "Attributes regrouping in Fuzzy Rule Based Classification Systems : an intra-classes approach," in *IEEE/ACS 15th International Conference on Computer Systems and Applications (AICCSA)*, 2018, pp. 1–7, doi: 10.1109/AICCSA.2018.8612802.
- [3] F. Liu, A. Ahmed, S. Chai, Q. Geok, S. Ng, and D. K. Prasad, "RS-HeRR : a rough set-based Hebbian rule reduction neuro-fuzzy system," *Neural Computing and Applications*, vol. 2, pp. 1–15, 2020, doi: 10.1007/s00521-020-04997-2.
- [4] L. Dutu, G. Mauris, and P. Bolon, "A Fast and Accurate Rule-Base Generation Method for Mamdani Fuzzy Systems," in *IEEE Transactions on Fuzzy Systems, Institute of Electrical and Electronics Engineers*, pp. 715–733, 2018, doi: 10.1109/TFUZZ.2017.2688349.
- [5] K. Nozaki, H. Tanaka, and H. Ishibuchi, "Distributed representation of fuzzy rules and its application to pattern classification," *Fuzzy Sets and Systems*, vol. 52, no. 1, pp. 21–32, 1992, doi: 10.1016/0165-0114(92)90032-Y.
- [6] R. Chen and Y. Hu, "A Novel Method for Discovering Fuzzy Sequential Patterns Using the Simple Fuzzy Partition Method," *Journal of The American Society for Information Science and Technology*, vol. 54, no. 7, pp. 660–670, 2003, doi:10.1002/asi.10258.
- [7] T. Chen, Q. Shen, P. Su, and C. Shang, "Refinement of Fuzzy Rule Weights with Particle Swarm Optimisation," in *14th UK Workshop on Computational Intelligence (UKCI)*, 2014, pp. 1–7, doi: 10.1007/s00500-015-1922-z.
- [8] O. S. Sitompul, E. B. Nababan, and Z. Alim, "Adaptive Distributed Grid-Partition in Generating Fuzzy Rules," in *International Conference on Information & Communication Technology an System (ICTS)*, 2017, pp. 119–124, doi: 10.1109/ICTS.2017.8265656.
- [9] M. Mao, Q. Chen, and J. Sun, "Construction and Optimization of Fuzzy Rule-Based Classifier with a Swarm Intelligent Algorithm," *Mathematical Problems in Engineering*, pp. 1–12, 2020, doi: 10.1155/2020/9319364.
- [10] P. B. Manesh, K. Shahbazi, and S. Shahryari, "Application of Grid partitioning based Fuzzy inference system and ANFIS as novel approach for modeling of Athabasca bitumen and tetradecane mixture viscosity," *Petroleum Science and Technology*, vol. 37, no. 14, pp. 1613–1619, 2019, doi: 10.1080/10916466.2018.1471488.
- [11] M. Kumar and N. Yadav, "Fuzzy Rough Sets and Its Application in Data Mining Field," *Advances in Computer Science and Information Technology (ACSIT)*, vol. 2, no. 3, pp. 237–240, 2015.
- [12] C. Zhang, "Classification Rule Mining Algorithm Combining Intuitionistic Fuzzy Rough Sets and Genetic Algorithm," *International Journal of Fuzzy Systems*, vol. 22, pp. 1694–1715, 2020, doi: 10.1007/s40815-020-00849-2.
- [13] K. P. Lin, K. C. Hung, and C. L. Lin, "Rule Generation Based on Novel Kernel Intuitionistic Fuzzy Rough Set Model," *IEEE Access*, vol. 6, pp. 11953–11958, 2018, doi: 10.1109/ACCESS.2018.2809456.
- [14] H. Slim and S. Nadeau, "A Mixed Rough Sets / Fuzzy Logic Approach for Modelling Systemic Performance Variability with FRAM," *Sustainability*, vol. 12, no. 5, pp. 1–21, 2020, doi: 10.3390/su12051918.
- [15] M. Landowski and A. Landowska, "Usage of the rough set theory for generating decision rules of number of traffic vehicles," *Transportation Research Procedia*, vol. 39, no. 2018, pp. 260–269, 2019, doi: 10.1016/j.trpro.2019.06.028.
- [16] J. Pant, R. P. Pant, A. Juyal, H. Pant, "Rule Generation Methods for Medical Data (Heart Disease) Using Rough Set Approach," in *International Conference on Advances in Engineering Science Management & Technology*, 2019, pp.1–12.
- [17] X. Yang, "Indiscernibility Relation, Rough Sets and Information System," in *Incomplete Information System and Rough Set Theory*, pp. 3–42, 2012, doi: 10.1007/978-3-642-25935-7_1.
- [18] Z. Pawlak, "Some issues on rough sets," in *Transactions on Rough Sets I*, pp 1–58, 2004, doi: 10.1007/978-3-540-27794-1_1.
- [19] A. Janusz, Cbergmeir, D. Jankowski, and A. Fiedukowicz, *RoughSets-package*. Retrieved from <https://github.com/janusza/RoughSets>, 2022.
- [20] R. Vashist and M. Garg "Rule Generation based on Reduct and Core: A Rough Set Approach," in *International Journal of Computer Applications*, vol. 29, no. 9, pp. 1–5, 2011, doi: 10.5120/3595-4989.
- [21] M. Michalak, M. Dubiel, and J. Urbanek, "Aplication For Logical Expresion Processing," *Computer Science & Information Technology (CS & IT)*, pp. 1–9, 2016, doi: 10.5121/csit.2016.60801.
- [22] I. A. Kesuma, M. Zarlis, and E. B. Nababan, "Feature Selection and Rule Extraction Based on Variable Precision Rough Set," in *The 3rd International Conference on Computing and Applied Informatics*, 2019, vol. 1235, no. 1, pp. 1–7, doi: 10.1088/1742-6596/1235/1/012052.
- [23] J. Kacprzyk and W. Pedrycz, "Soft computing in database and information management," in *Springer Handbook of Computational Intelligence*, pp. 295–312, 2015, doi:10.1007/978-3-662-43505-2.
- [24] J. G. Marín-Blázquez and Q. Shen, "From approximative to descriptive fuzzy classifiers," *IEEE Transactions on Fuzzy Systems*, vol. 10, no. 4, pp. 484–497, 2002, doi: 10.1109/TFUZZ.2002.800687.
- [25] H. Ishibuchi, K. Nozaki, and H. Tanaka, "Efficient fuzzy partition of pattern space for classification problems," in *Fuzzy Sets and Systems*, vol. 59, no. 3, pp. 295–304, 1993, doi: 10.1016/0165-0114(93)90474-V.

BIOGRAPHIES OF AUTHORS

Murni Marbun    got master from Universiti STMIK Eresha in 2014 and graduate from Universitas Sumatera Utara in 1998. She is a lecturer in Department of Information Technology, STMIK Pelita Nusantara, Medan in 2014. His research interest includes computer science, fuzzy logic, and decision support systems. She can be contacted at email: murni.marbun@students.usu.ac.id.



Opim Salim Sitompul    received the Ph.D. degree in information science from Universiti Kebangsaan Malaysia, Selangor, in 2005. He is currently a Professor with the Department of Information Technology, Sumatera Utara University, Medan, Indonesia. His skill and expertise are in AI, data warehousing, and data science. He can be contacted at email: opim@usu.ac.id.



Erna Budhiarti Nababan    completed her Ph.D. on Science and Management Systems from Universiti Kebangsaan Malaysia, Selangor in 2010. She is currently a senior lecturer at Department of Information Technology, Universitas Sumatera Utara, Medan, Indonesia. She can be contacted at email: ernabrn@usu.ac.id.



Poltak Sihombing    received the Ph.D. degree in computer science from Universiti Sains Malaysia, Malaysia, in 2010. He is currently a Professor with the Department of Information Technology, Sumatera Utara University, Medan, Indonesia. His skill and expertise are in microcontroller and IoT, intelligent system, and information retrieval. He can be contacted at email: poltak@usu.ac.id.