

ASSESSMENT MEASURES OF AN ENSEMBLE CLASSIFIER BASED ON THE DISTRIBUTIVITY EQUATION TO PREDICT THE PRESENCE OF SEVERE CORONARY ARTERY DISEASE

EWA RAK ^{a,*}, ADAM SZCZUR ^b, JAN G. BAZAN ^b, STANISŁAWA BAZAN-SOCHA ^c

^aInstitute of Mathematics
University of Rzeszów
Pigonia 1, 35-310 Rzeszów, Poland
e-mail: erak@ur.edu.pl

^b Institute of Computer Science
University of Rzeszów
Pigonia 1, 35-310 Rzeszów, Poland
e-mail: {aszczur, jbazan}@ur.edu.pl

^cDepartment of Internal Medicine
Jagiellonian University Medical College
Skawinska 8, 31-066 Kraków, Poland
e-mail: stanislawa.bazan-socha@uj.edu.pl

The aim of this study is to apply and evaluate the usefulness of the hybrid classifier to predict the presence of serious coronary artery disease based on clinical data and 24-hour Holter ECG monitoring. Our approach relies on an ensemble classifier applying the distributivity equation aggregating base classifiers accordingly. Such a method may be helpful for physicians in the management of patients with coronary artery disease, in particular in the face of limited access to invasive diagnostic tests, i.e., coronary angiography, or in the case of contraindications to its performance. The paper includes results of experiments performed on medical data obtained from the Department of Internal Medicine, Jagiellonian University Medical College, Kraków, Poland. The data set contains clinical data, data from Holter ECG (24-hour ECG monitoring), and coronary angiography. A leave-one-out cross-validation technique is used for the performance evaluation of the classifiers on a data set using the WEKA (Waikato Environment for Knowledge Analysis) tool. We present the results of comparing our hybrid algorithm created from aggregation with the distributive equation of selected classification algorithms (multilayer perceptron network, support vector machine, k -nearest neighbors, naïve Bayes, and random forests) with themselves on raw data.

Keywords: ensemble method, distributivity equation, aggregation function, accuracy, precision, sensitivity, CAD, Holter ECG.

1. Introduction

Machine learning (ML), as a branch of artificial intelligence (AI) that attempts to imitate intelligent behavior, is one of the most promising approaches to solving difficult decision-making problems. The general idea is very simple: instead of modeling a solution explicitly, a domain expert provides example data that demonstrate the desired behavior on representative

problem instances. The appropriate ML algorithm is then trained on these examples to best reproduce the expert's solutions and generalize them to new, unobserved data. As the field of ML develops, we are able to train the computer to solve increasingly complex tasks. Moreover, the higher and higher obstacles constantly stimulate further development. One method of improving “what we already have” is combined (known also as a hybrid, multiple, or ensemble) classification. Ensemble (hybrid) methods are known as learning algorithms that train a set of classifiers

*Corresponding author

and combine them to achieve the best prediction measures (Dietterich, 2000).

Various methods proposed over the years use various strategies for computing this combination (Kaczmarek-Majer and Kiersztyn, 2022; Kunapuli, 2023). The most fundamental concepts of ensemble methods consist of two main stages, which are the production of multiple base classifier models and their combination via aggregation. Aggregation functions proved to be an effective tool in many application areas (Beliakov *et al.*, 2016). It simply refers to calculations performed on a data set to get a single number that accurately represents the underlying data. There are also many approaches using domain knowledge and improving the quality of data mining models (e.g., Garg, 2021). Thus, the use of aggregation functions can be treated as a way to use domain knowledge to improve the quality of classifiers.

A skillful selection of both, ensemble classifiers and aggregation, can bring many benefits. Such advantages are noted in our recently developed approach that included the ensemble (hybrid) method applying the distributivity law (Rak *et al.*, 2020; Rak and Szczur, 2021). Experiments were performed on the cyber-attacks in the military network data set obtained from the machine learning repository UCI (Dua and Casey, 2019). Another goal we set for ourselves was to check whether this method also translates into medical data. Thus, this paper continues research with a novel hybrid approach to increase the main classification measures: accuracy, sensitivity, and precision while minimizing base classifiers by applying the distributivity law that aggregates classifiers appropriately.

This paper includes results of experiments that have been performed on medical data (clinical data and Holter electrocardiogram (ECG) monitoring records). From the medical point of view, the study involves the prediction of coronary arteriosclerosis presence in patients with stable angina. Coronary artery disease (CAD) (also known as coronary heart disease (CHD), coronary microvascular disease (CMD), cardiovascular disease (CVD) or atherosclerosis, arrhythmia and arterial thrombosis) is the most common type of heart disease. It is the leading cause of death in all countries for both men and women. As reported in 2021 by the World Health Organization (WHO), cardiovascular diseases are the cause of death for 17.9 million patients each year, which represents nearly 32% of all deaths worldwide (see <https://www.who.int/health-topics/cardiovascular-diseases>).

Diagnosis of CAD is made using various tests such as an electrocardiogram or a stress test. Treatment for CAD includes lifestyle changes, medications, and sometimes, cardiac procedures or surgery—coronary angiography. Prevention consists of modifying reversible

risk factors (e.g., hypercholesterolemia, hypertension, physical inactivity, obesity, diabetes, smoking). We propose the use of clinical data together with Holter electrocardiogram recordings as prospective candidate data for coronary artery stenosis prediction. The proposed ensemble method with the use of the distributivity law helps us to determine the management of patients with stable angina, including the need for coronary intervention, without performing invasive diagnostic procedures such as angiography. To some extent, it also works as a screening tool for all patients with CAD. Similar considerations for other classification methods were conducted by Bazan *et al.* (2020).

The comparison of classifiers and using the most predictive classifier is very important. Each of the classification methods shows different efficiency and accuracy based on the kind of data sets (Kim, 2008). Five different classification algorithms from WEKA API (Frank *et al.*, 2016) applying the distributivity equation which aggregates the classifiers accordingly were used, and their quality was compared to the prediction measures based on the confusion matrix obtained on raw data. Basic algorithms for the custom aggregation method that have been investigated are multilayer perceptron network (MLP), support vector machine (SVM), k -nearest neighbor (kNN), naïve Bayes (NB) and random forests (RFs).

The rest of the paper is organized as follows. Section 2 provides an overview of published articles on CAD classification using various ML and data mining approaches. In Section 3, notions connected with aggregation functions and distributivity between them are recalled. Section 4 explains the used data set. Section 5 describes the experimental setting, and evaluates the results of applying the ensemble classification method based on the distributivity law (with five selected classifiers) on the medical data set through the classification measures for each case of coronary artery disease. Finally, Section 6 summarizes this work.

2. Review of the literature

Modern methods of data analysis are provided by modern multivariate statistics, where classification methods are of particular practical importance. Supervised learning provides a powerful tool to classify and process data using machine language. A classification algorithm is a procedure for selecting a hypothesis from a set of alternatives that best fit a set of observations. Classifier algorithms are widely utilized in data mining. They can generate a solid prediction model based on a set of features during the training phase. These features belong to certain labeled classes. The generated prediction model is used later to predict new classes.

There are many different classification techniques

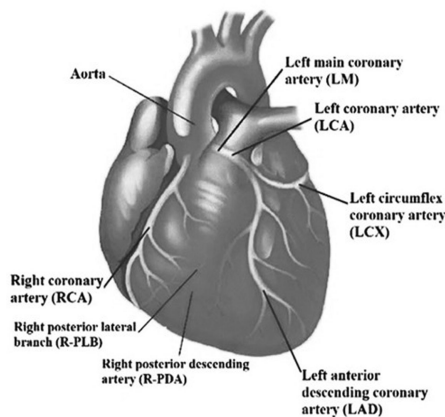


Fig. 1. Model of the human heart.

including support vector machines (SVMs) (Tanveer *et al.*, 2019), artificial neural networks (ANNs) (Schmidhuber, 2015), decision trees (DTs) (Jaworski *et al.*, 2018), the multilayer perceptron (MLP) (Du and Swamy, 2014) or *k*-nearest neighbors (kNN) (Zhang *et al.*, 2018). Each of these techniques has its strengths and weaknesses. These methods have been widely investigated in broad areas such as medical applications (e.g., breast cancer (Kowal *et al.*, 2021), Parkinson's disease (Bernardo *et al.*, 2021) or electrocardiogram (ECG) signals of the human heart (Patro *et al.*, 2022)) for purposes such as screening, risk stratification, prediction, and assisted decision-making (Castaneda *et al.*, 2015).

The mortality rate among various diseases is an important factor in undertaking these studies. According to WHO, the leading cause of death is coronary artery disease (CAD). It occurs when there is an obstruction (of more than 50%) in at least one of the coronary arteries (Zipes *et al.*, 2018). There are three major arteries of the heart: left anterior descending artery (LAD), left circumflex artery (LCX), and right coronary artery (RCA) (see Fig. 1). Thus, early detection of CAD is critical to avoid further increases in the risk. Coronary angiography is required to conclusively diagnose CAD. However, it is invasive and may lead to various complications, such as artery dissection, arrhythmia, and even death. Moreover, image-based detection techniques are costly and not applicable for screening large populations, especially in developing countries. Due to these shortcomings and the life-threatening nature of angiography, researchers have been continuously looking for noninvasive, economical, fast, and reliable techniques for early detection of CAD. ML algorithms are techniques used for this purpose (Krittanawong *et al.*, 2020; Alfai *et al.*, 2022). There are works that were published 30 years ago (see Akay, 1992).

Several indicators have been used in the literature, including accuracy, sensitivity, specificity and the f-score for model evaluation. However, the overall performance

of the model depends on two key factors: the data set and the choice of the ML method. CAD detection mainly uses supervised classification algorithms for feature processing and decision making. The most frequently used ML methods for CAD detection: are ANN, DTs, and SVM. They are the most common methods, which have been applied to almost all data sets that have been reported in the literature to date (see the current review paper by Alizadehsani *et al.* (2019)). Moreover, the most common assessment metric for CAD detection is still only accuracy. Relying solely on accuracy can be confusing, especially for highly unbalanced data sets. Therefore, accuracy should be always evaluated and interpreted in conjunction with other metrics, including sensitivity and precision.

We are interested in studies on three stenoses of LAD, LCX, and RCA arteries. The best method for diagnosing CAD is angiography, in which three main outputs are determined: (i) whether or not the patient has CAD (ii) which artery is stenotic, and (iii) the percentage of the stenosis. The first output is already well considered in the literature. In contrast, the second output is only studied in few articles (Alizadehsani *et al.*, 2013; 2016; Babaoglu *et al.*, 2009), with the highest accuracy of 86.1% for LAD stenosis diagnosis (Alizadehsani *et al.*, 2016). Surprisingly, there is no study that reports on the third output. Hence, reducing or even eliminating the use of angiography due to its costs and side effects requires extensive research and work on the second and third outputs of angiography. When comparing the performance of the algorithms described by Alizadehsani *et al.* (2019), "naïve Bayes was unable to exceed C4.5. Additionally, naïve Bayes fared worse than SVM in all but one. The ANN performance in all but one article was worse than that of the SVM or fuzzy rule-based system (FRBS)."

3. Distributivity of aggregation functions

The aggregation process is a synthesis of many numerical data to a single value, in some way, a representative for all of them. This type of projection method of multidimensional space input data to one dimension is usually carried out by the so-called aggregation functions (also known as aggregation operators). They were formalized forty years ago (Dombi, 1982), and since then have been extensively investigated (Grabisch *et al.*, 2009; Beliakov *et al.*, 2016). It is worth mentioning that aggregation functions have been successfully applied to solve many pragmatic application problems including decision making, hierarchical information fusion, classification, image processing, fuzzy and control systems, etc.

The choice of the aggregation function should be based upon properties dedicated by the framework in

which the aggregation is performed. Of the huge number of classes of aggregation functions, averaging functions (means) (see, e.g., Beliakov *et al.*, 2016, p. 55) and triangular norms (see Klement *et al.*, 2000, p. 6) are usually easily applied to the classification problem (see Table 1).

Distributivity specifies the relationship between two binary functions, including aggregation functions.

Definition 1. (Aczél, 1966, p. 318) Let F and G be some binary functions in a non-empty set U , where F is symmetric. We say that F is distributive over G if for all $X, Y, Z \in U$ the following equality is fulfilled:

$$F(X, G(Y, Z)) = G(F(X, Y), F(X, Z)). \quad (1)$$

The lack of distributivity is a big problem in any algebraic transformations, and therefore also in computer modeling. In general, aggregations are not distributive from each other. The sufficient condition under which one aggregation function is distributive with respect to another is the idempotency of the second aggregation. However, the sufficient condition(s) are no longer easy to indicate, since they are different for various classes of aggregation functions and strictly depend on the structures of these functions. For t-norms, t-conorms, and means, that can be easily applied to classification and decision-making problems, we rely on the results of Drewniak *et al.* (2008, Lemmas 1 and 2) as well as Drewniak and Rak (2010, Tables 3 and 4). Out of all pairs of aggregation functions known so far for which the distributivity equation (1) formally holds, we have selected (see Table 1) and used in the elaborated algorithm the following equalities:

- D1 $T_P(X, M_\wedge(Y, Z)) = M_\wedge(T_P(X, Y), T_P(X, Z))$,
- D2 $T_P(X, M_\vee(Y, Z)) = M_\vee(T_P(X, Y), T_P(X, Z))$,
- D3 $T_P(X, M_A(Y, Z)) = M_A(T_P(X, Y), T_P(X, Z))$,
- D4 $M_A(X, M_A(Y, Z)) = M_A(M_A(X, Y), M_A(X, Z))$,
- D5 $T_P(X, M_H(Y, Z)) = M_H(T_P(X, Y), T_P(X, Z))$,
- D6 $M_P(X, M_P(Y, Z)) = M_P(M_P(X, Y), M_P(X, Z))$,
- D7 $T_E(X, M_\wedge(Y, Z)) = M_\wedge(T_E(X, Y), T_E(X, Z))$,
- D8 $T_E(X, M_\vee(Y, Z)) = M_\vee(T_E(X, Y), T_E(X, Z))$,
- D9 $T_H(X, M_\wedge(Y, Z)) = M_\wedge(T_H(X, Y), T_H(X, Z))$,
- D10 $T_H(X, M_\vee(Y, Z)) = M_\vee(T_H(X, Y), T_H(X, Z))$.

4. Data set

4.1. Medical background. Our approach is illustrated using data that represent the medical treatment of patients with stable coronary artery disease, which is a major health problem worldwide and is one of the leading causes

of high mortality rates in industrialized countries. It is called angina, due to one of its main symptoms—chest pain, arising from ischemia of the heart muscle. The main cause of CAD is artery stenosis and the consequences of CHD depend largely on the number, degree, and localization of artery stenosis.

The data set, collected by the Department of Internal Medicine, Collegium Medicum, Jagiellonian University, Kraków, Poland, relates to 152 patients subjected to elective coronary angiography with possible percutaneous angioplasty. The data set contains clinical data (some of them are shown in Table 2), data from Holter ECG (24-hour ECG monitoring), and coronary angiography. Holter ECG study was carried out within 24 hours before the angiography, which is the current diagnostic standard of anatomic coronary vessel evaluation which permits the determination of the therapeutic plan and prognosis. In the case of an unaltered coronary flow, pharmacological treatment is applied otherwise, revascularization is also needed. However, coronary angiography (coronagraphy) is a very sensitive method and has its limitations. As an invasive investigation, it is relatively expensive, and it carries risks including a mortality rate of approximately 1 in 2000.

It would not be appropriate or practical to perform invasive investigations on all patients with a coronary heart disease diagnosis. Given the high incidence and prevalence of CAD, a non-invasive test to reliably assess the coronary arteries would be clinically desirable. Our experiments are therefore an attempt to use medical data (clinical data together with electrocardiographic Holter recordings), to construct a classifier that allows identifying which patients with CAD need revascularization surgery, and for whom pharmacological treatment is sufficient because their stenosis is small. This diagnosis is made without coronary angiography. However, coronary angiography was performed for all patients in the data, as its results were used to define the binary decision attribute (see results collected in Table 3).

4.2. Experimental data. The experimental data contain an ECG recorded using the Holter method, enriched with clinical data of patients with stable ischemic heart disease with sinus rhythm in the ECG recording. Patients were recruited from among those admitted to the Department of Internal Medicine and Heart Diseases for elective surgery of coronary angiography with possible angioplasty and stent implantation. Immediately after coronary angiography, results were subjected to angiographic analysis, which enabled patients to be qualified for percutaneous treatment. In patients qualified for the above-mentioned treatment, one-time coronary angioplasty with or without stent implantation was performed. Before and after surgery, 24-hour Holter

Table 1. Means and t-norms considered in this research.

Mean	Name
$M_{\wedge}(X, Y) = \min(X, Y)$	minimum
$M_{\vee}(X, Y) = \max(X, Y)$	maximum
$M_A(X, Y) = \frac{X+Y}{2}$	arithmetic mean
$M_H(X, Y) = \begin{cases} 0 & \text{if } X = Y = 0 \\ \frac{2XY}{X+Y}, & \text{otherwise} \end{cases}$	harmonic mean
$M_P(X, Y) = \sqrt{\frac{X^2+Y^2}{2}}$	power mean
T-norm	Name
$T_P(X, Y) = X \cdot Y$	product t-norm
$T_E(X, Y) = \frac{X \cdot Y}{2 - (X + Y - X \cdot Y)}$	Einstein t-norm
$T_H(X, Y) = \begin{cases} 0 & \text{if } X = Y = 0 \\ \frac{XY}{X+Y-X \cdot Y} & \text{otherwise} \end{cases}$	Hamacher t-norm

Table 2. Short clinical characteristics of patients.

Feature	Value
Number of patients	$N = 152$ (100%)
Age	40–87 (68.75 $\pm$ 9.643)
Sex (M/F)	90 / 62 (59% / 41%)
Hyperlipidemia (Y/N)	44 / 108 (29% / 71%)
Obesity (Y/N)	57 / 95 (37.5% / 62.5%)
Smoking tobacco (Y/N)	44 / 108 (29% / 71%)

ECG monitoring was performed on all patients. Each time, after completion of the examination, the data saved on the portable memory card of the recording set were loaded into the memory of a desktop computer and then subjected to an automatic analysis using the software provided by the manufacturer of the ECG recorder. These data include HOLTER collection, which contains data from 152 patients collected in 2015 and 2016 using the 12-channel R12 recorder of the BTL CardioPoint-Holter H600 v2-23 system. Angiographic data provide detailed information on the percent stenosis for each of the assessed coronary angiographies. Patients were selected for the collection without complex cardiac arrhythmias, such as supraventricular or ventricular extrasystoles, which prevent proper ECG analysis.

Only medical data with Holter electrocardiographic records before coronary angiography, supported by clinical data, were used for the experiments. In particular, it includes an accurate description of the clinical condition of patients (age, sex, medical diagnosis), comorbidities, pharmacological treatment, laboratory test results (including troponin level, CRP protein, cholesterol, LDL), and many Holter parameters concerning the number and type of arrhythmias, changes in the PQ interval, changes in the ST segment or heart rate variability (HRV) in the domain of time and frequency,

Table 3. Angiographic characteristics of patients.

	Holter
Result of coronary angiography	$N = 152$ (100%)
No significant stenosis in the coronary arteries	86 (56.5%)
1 stenosis	32 (21%)
2 stenoses	19 (12.5%)
3 stenoses	15 (10%)

and changes in the QT interval. The 24-hour Holter recording for each patient was aggregated into a single row in a data table. In the aggregated data, 7 new features were defined for each of the original attributes. The new features were based on statistical measures like first and last values, minimum, maximum, mean, standard deviation and total. The values of these attributes were calculated for each object (patient) based on the values of their time points. Finally, the collection contains 595 attributes.

The obtained data were loaded into a database using an importer created in the Java environment. The data were preprocessed—textual (symbolic) features were omitted and individual patient data were merged. The original data set has some imbalances that may affect the classification accuracy of the algorithm. Therefore, it was necessary to balance the original data set.

This was accomplished with the use of the filter `weka.filters.supervised.instance.Resample-B1.0-S1-Z100.0`, where $(86 + 32 + 19 + 15)/4 = 38$ was the average number of objects per class. Next, the data were exported to a text file.

5. Classifier aggregation method based on the distributivity law

We recall the novel ensemble approach proposed in our previous paper (Rak *et al.*, 2020) to increase the accuracy of a classification and, at the same time, minimize a group of base classifiers by applying the distributivity law which will aggregate classifiers accordingly (see the scheme in Fig. 2).

The proposed algorithm was subjected to further modifications in order to improve the accuracy of the classification, as well as sensitivity, false positive rate (fall-out), and precision. These are the so-called classification quality measures based on a confusion matrix.

The confusion matrix is a situation analysis table for summarizing the prediction results of classification models in ML. Here ‘actual class’ is also known as ground truth (GT) value and ‘predicted class’ is the output of the model. As shown in Table 4, TP is the number of correctly predicted positive cases, FN is the number of incorrectly predicted positive cases, FP is the number of incorrectly predicted negative cases, and TN is the number of correctly predicted negative cases. Our experimental evaluation measures based on a confusion matrix are the following:

- accuracy

$$ACC = \frac{TP + TN}{TP + TN + FP + FN}$$

is one of the main model assessment parameters that define the proportion of correct classifications;

- sensitivity (recall or true positive rate)

$$TPR = \frac{TP}{TP + FN}$$

measures the proportion of correct CAD predictions to all cases that have CAD;

- false positive rate

$$FPR = \frac{FP}{FP + TN}$$

can be defined as the percentage of healthy people that the classification model incorrectly identifies them as not having CAD;

- precision (positive predictive value)

$$PPV = \frac{TP}{TP + FP}$$

is a measure of how reliable positive predictions are, i.e., the percentage of positive predictions that are actually positive (the percentage of people with a positive test result, in whom the diagnosis was significantly confirmed).

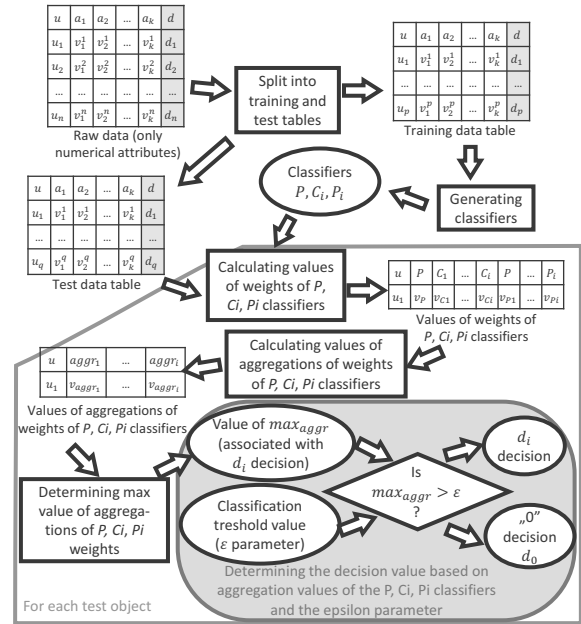


Fig. 2. Scheme of the proposed approach.

Table 4. Confusion matrix in predicting the presence of CAD.

Actual ↓ Predicted →	Positive	Negative
Positive	TP	FN
Negative	FP	TN

5.1. Steps of the proposed ensemble method. Let d_0 be the decision class "0" (no significant stenosis in the coronary arteries) and d_1, d_2, d_3 denote 3 degrees of vascular disease, respectively. Moreover, $\varepsilon \in (0, 1)$ denotes the threshold for classifying an object into one of the d_i classes, where $(i = 1, 2, 3)$. The use of ε parameter in experiments is described later in this chapter, where the algorithm for the classification of the test object has been placed.

The process of building classifiers and determining decision values for test objects is given by Algorithm 1.

5.2. Experiments. The algorithm, along with the use of five classification methods, was implemented and tested in Java programming language (using the Weka API library for construction of classifiers P, C_i, P_i).

The following classifiers were used for the experiments:

- kNN (IBk) with Euclidean distance and $k = \sqrt{n}$, where n is the number of objects,
- MLP with default options (500 epochs, seed = 0, hidden layers = (attribs + classes)/2),

Algorithm 1. Classifier aggregation based on the distributivity law (CADL).

Input: Training and test data tables (TR and TST, respectively)

Output: Decision value(s) for tested object(s)

Step 1. Construct the classifiers on TR data as follows:

- (i) Classifier P – binary classifier used to pre-classify an object as to whether it belongs to decision class d_0 (marked as class ‘a’) or to some other decision class $\neg d_0 = \{d_1, d_2, d_3\}$ (marked as class ‘b’);
- (ii) Collection of binary classifiers C_i ($i = 1, 2, 3$) used to distinguish objects with decision value d_i (marked as class ‘b’) from objects from other classes $\neg d_i = \{d_1, d_2, d_3\} \setminus \{d_i\}$ (marked as a class ‘a’) without the class d_0 ;
- (iii) Collection of binary classifiers P_i ($i = 1, 2, 3$) used to distinguish objects with decision value d_i (marked as class ‘b’) from objects of the class d_0 (marked as class ‘a’).

For each object u_j from table TST, follow Steps 2–6:

Step 2. Determine the probability of membership u_j to decision class with label ‘b’ (weight for a class marked ‘b’). Next, construct the weight table of classifiers P, C_i, P_i for the class marked ‘b’ (classifiers as columns, u_j object as row, weights (probabilities) for class ‘b’ as values).

Step 3. For a given distributivity law-hand satisfying $D1$ – $D10$ compute the value of its left-hand side L (as $L = P$) based on weights from Step 2 and considering obtained classifiers X, Y, Z in the setting $X = P, Y \in C_i, Z \in P_i$.

Step 4. Fix the parameter $\varepsilon \in (0, 1)$ (the same for each tested instance u_j).

Step 5. Determine the maximum for previously calculated values of the left-hand side of the distributivity law ($\max_{\text{aggr}} = \max(D1_{(X,Y,Z)}, \dots, D10_{(X,Y,Z)})$).

Step 6. Propose a decision value for the object u_j as follows: If the maximum value of the left-hand side ($\max_{\text{aggr}}$) of distributivity was obtained for the i -th decision (after aggregating the weights of the classifiers P, C_i, P_i) and $\max_{\text{aggr}} > \varepsilon$, then propose the i -th decision value for the given object (respectively 1, 2 or 3 stenoses); otherwise, ($\max_{\text{aggr}} \leq \varepsilon$) propose a neutral decision – the lack of stenoses (class ‘0’).

- RFs with default options (100 iterations (trees), attributes to randomly investigate = $\log_2(\text{predictors}) + 1$, seed = 1).

It uses the above selected examples of pairs of distributive aggregation functions denoted as $D1$ – $D10$. We used leave-one-out cross-validation (LOO CV) to test the quality of classifiers. It is usually employed when the size of a given data set is small. The LOO CV technique involves a single object from the original data set as validation data and the remaining observations as training data. This is repeated in such a way that each observation in the sample is used once as validation data. Experiments are performed using 582 numerical features. The following parameters were used as measures of the success (or failure) of the classification: accuracy, accuracy for positive examples (sensitivity), false-positive rate, and precision for positive examples.

The original data set has a certain imbalance. Although we have balanced data, accuracy itself cannot measure the effective performance of the model. A comparison of selected results of TPR, FPR, PPV and overall ACC measurements, on the original raw data and the newly created data using the classifier aggregation method with five algorithms (SVM, MLP, kNN, NB, and RFs) with a fixed $\varepsilon \in (0, 1)$, due to their size, is presented in Tables A1–A11 in Appendix. In turn, the file constituting the full version of the results can be found on the GitHub platform: https://github.com/Ama79/Results_AMCS.

In addition, selected results, covered by the aforementioned tables, are also presented in a graphical form (see Figs. A1–A11 in Appendix).

5.3. Discussion of results. Since data sets usually vary in the number of samples and features, it is not easy to compare ML techniques in terms of performance. Nevertheless, each new CAD detection method brings us closer to solving the decision problem of whether further in-depth (expensive and dangerous) diagnostics is really necessary. The overall performance results of our method for the difficult medical data set under consideration are very promising.

Depending on the standard classification method chosen (SVM, MLP, kNN, NB, RFs) and the adopted distributivity equation, the experimental results obtained from our method for predicting coronary stenosis in coronary artery disease differ. Thus, it is difficult to clearly determine the suitability of exactly one of the selected standard classifiers for the hybrid method, which seeks to obtain the best possible measures of its evaluation in relation to raw data. It can certainly be seen that in the cases of $P = kNN, C_i = kNN, P_i = kNN$ and $P = NB, C_i = NB, P_i = NB$ there are no satisfactory results. The use of RFs, in turn, yields good results,

- SVM (SMO) with default options (PolyKernel, logistic calibrator),
- NB with default options (normal distribution for numeric attributes, no discretization),

however, similar to those on raw data, with improvements in sensitivity (from 86.8% to 100%), precision (from 82.5% to 100%), and overall accuracy (from 85.5% to 86.8%) in predicting the lack of arterial stenoses using equation D1. Similar results, but covering globally all cases of coronary artery stenosis prediction, are obtained employing MLP and SVM with equation D8 and a fixed $\varepsilon = 0.7$. However, the best results of the proposed method are when different classifiers are combined, especially when $P = \text{MLP}$, $C_i = \text{SVM}$, $P_i = k\text{NN}$ using equations D6 (with ε from 0.7 to 0.8) and D9 and D10 (with $\varepsilon = 0.9$).

This method in relation to raw data better identifies patients who do not have stenosis in particular, yields it 100% precision, 100% sensitivity, and 0% false positive rate. In turn, for patients with significant coronary artery stenosis, the method used offers the following best measures for assessing the diagnosis:

- 1 stenosis: 88.9% precision, 73.7% sensitivity and 1.8% false positive rate;
- 2 stenoses: 87.2% precision, 92.1% sensitivity and 3.5% false positive rate;
- 3 stenoses: 87.5% precision, 92.1% sensitivity and 4.4% false positive rate; the overall accuracy of the proposed model is 86.8%.

It can be certainly stated that the best measures were also influenced by the selection of aggregation functions satisfying the distributivity equation (1). Generally, the best results are for equations, in which there is a combination of the product t-norm, Einstein and Hamacher t-norm with maximum. Results obtained justify our approach of using the idea of distributivity of aggregation functions in the classification method. Adapting this hybrid classification method with the simultaneous use of, for example, MLP and SVM with RFs instead of kNN (currently in our research phase) in highly desirable medical decision support can bring many benefits for effective diagnosis.

6. Conclusions

In machine learning a combination of classifiers (known as ensemble or hybrid classification) often outperforms individual ones. While many ensemble approaches exist, it remains, however, a difficult task to find a suitable ensemble configuration for a particular data set. Within our research interests are multi-class datasets, and for them, we are trying to develop a reasonably universal algorithm using nonstandard mathematical tools. The aim of this paper was to apply the ensemble method (classifiers aggregation based on the distributivity law) to predict the presence of serious coronary artery disease based on clinical and ECG data. It presents a comparative

assessment for a novel ensemble construction method that uses a variety of standard supervised classification algorithms and applies the distributivity law which aggregates these classifiers accordingly. The underlying algorithms for the custom aggregation method that were connected and compared with raw data by accuracy, precision, TPR, and FPR measures are the multilayer perceptron network, k -nearest neighbors, support vector machine, naive Bayes, and random forests.

Our experimental results suggest that the hybrid approach can generate ensembles that outperform traditional algorithms in terms of classification precision, sensitivity, and accuracy. With the appropriate value of the parameter ε , we obtained results up to 15% (TPR) 40% (PPV) and 2% (ACC) better (than for raw data) for individual decision classes. Thus, this approach can be useful for clinicians in the management of patients with coronary artery disease, in particular in the face of limited access to invasive diagnostic tests, i.e., coronary angiography or in the case of contraindications to its performance (allergy to contrast administered during coronary angiography, poor general condition of the patient, other acute illnesses). The proposed method is important for physicians who treat patients with coronary artery disease in their daily practice.

The classification problem raised in this paper is related to the challenge in cardiology published in 2003 (cf. Moody and Jager, 2003). At that time, unfortunately, there was no widespread feedback. The problem is still relevant and all attempts to support the diagnosis of heart diseases with ML methods are very desirable. Therefore, we would like to mention that we have made some additions to the knowledge in support of CAD diagnosis, meeting both the expectations and challenges of cardiology and the recommendation of Alizadehsani *et al.* (2019): "... Only five studies reported the application of ML in determining which artery is stenotic and the classification results for these studies are poor. As we discussed above, the highest accuracy in identifying CAD stenosis was 86.1% (Alizadehsani *et al.*, 2016). Hence, improvement of the algorithms in this area is required to realize better prediction performance. Therefore, additional research is strongly recommended. Moreover, there is no study in which the category three output (percentage of stenosis) is determined via ML methods." In future work, we will add different attribute selection techniques to the current method or supplement them where they are missing. In addition, we plan to test more data sets to be able to get an even more accurate evaluation.

References

- Aczél, J. (1966). *Lectures on Functional Equations and Their Applications*, Academic Press, New York/London.

- Akay, M. (1992). Noninvasive diagnosis of coronary artery disease using a neural network algorithm, *Biological Cybernetics* **67**(4): 361–367.
- Alfaidi, A., Aljuhani, R., Alshehri, B., Alwadei, H. and Sabbeh, S. (2022). Machine learning: Assisted cardiovascular diseases diagnosis, *International Journal of Advanced Computer Science and Applications* **13**(2): 135–141.
- Alizadehsani, R., Abdar, M., Roshanzamir, M., Khosravi, A., Kebria, P.M., Khozeimeh, F., Nahavandi, S., Sarrafzadegan, N., Acharya, U.R. and Abdar, M. (2019). Machine learning-based coronary artery disease diagnosis: A comprehensive review, *Computers in Biology and Medicine* **111**(4): 103346.
- Alizadehsani, R., Habibi, J., Alizadeh-Sani, Z., Mashayekhi, H., Boghrati, R., Ghandeharioun, A., Khozeimeh, F. and Alizadeh-Sani, F. (2013). Diagnosing coronary artery disease via data mining algorithms by considering laboratory and echocardiography features, *Research in Cardiovascular Medicine* **2**(3): 133–139.
- Alizadehsani, R., Zangoeei, M.H., Hosseini, M.J., Habibi, J., Khosravi, A., Roshanzamir, M., Khozeimeh, F., Sarrafzadegan, N. and Nahavandi, S. (2016). Coronary artery disease detection using computational intelligence methods, *Knowledge-Based Systems* **109**: 187–197.
- Babaoglu, I., Baykan, O.K., Aygul, N., Ozdemir, K. and Bayrak, M. (2009). Assessment of exercise stress testing with artificial neural network in determining coronary artery disease and predicting lesion localization, *Expert Systems with Applications* **36**(2): 2562–2566.
- Bazan, J.G., Bazan-Socha, S., Ochaba, M., Buregwa-Czuma, S., Nowakowski, T. and Woźniak, M. (2020). Effective construction of classifiers with the k -NN method supported by a concept ontology, *Knowledge and Information Systems* **62**: 1497–1510.
- Beliakov, G., Bustince, H. and Calvo, T. (2016). *A Practical Guide to Averaging Functions*, Springer, Cham.
- Bernardo, L.C., Damaševičius, R., de Albuquerque, V.H.C. and Maskeliūnas, R. (2021). A hybrid two-stage SqueezeNet and support vector machine system for Parkinson's disease detection based on handwritten spiral patterns, *International Journal of Applied Mathematics and Computer Science* **31**(4): 549–561, DOI: 10.34768/amcs-2021-0037.
- Castaneda, C., Nalley, K., Mannion, C., Bhattacharyya, P., Blake, P., Pecora, A., Goy, A. and Suh, K.S. (2015). Clinical decision support systems for improving diagnostic accuracy and achieving precision medicine, *Journal of Clinical Bioinformatics* **5**(4): 1–16.
- Dietterich, T.G. (2000). Ensemble methods in machine learning, in D. Haussler (Ed.), *Multiple Classifier Systems*, Springer, Heidelberg, pp. 1–12.
- Dombi, J. (1982). Basic concepts for the theory of evaluation: The aggregative operator, *European Journal of Operational Research* **10**(3): 282–293.
- Drewniak, J., Drygaś, P. and Rak, E. (2008). Distributivity equations for uninorms and nullnorms, *Fuzzy Sets and Systems* **159**(13): 1646–1657.
- Drewniak, J. and Rak, E. (2010). Subdistributivity and superdistributivity of binary operations, *Fuzzy Sets and Systems* **161**(2): 189–201.
- Du, K.-L. and Swamy, M.N.S. (2014). *Neural Networks and Statistical Learning*, Springer, London.
- Dua, D. and Casey, G. (2019). *UC Irvine Machine Learning Repository*, <http://archive.ics.uci.edu/ml/>.
- Frank, E., Hall, M.A. and Witten, I.H. (2016). *The WEKA Workbench—Online Appendix for Data Mining: Practical Machine Learning Tools and Techniques*, Morgan Kaufmann, Burlington.
- Garg, H. (2021). New exponential operation laws and operators for interval-valued q -rung orthopair fuzzy sets in group decision-making process, *Neural Computing and Applications* **33**: 13937–13963, DOI: 10.1007/s00521-021-06036-0.
- Grabisch, M., Marichal, J.L., Mesiar, R. and Pap, E. (2009). *Aggregation Functions*, Cambridge University Press, Cambridge.
- Jaworski, M., Duda, P. and Rutkowski, L. (2018). New splitting criteria for decision trees in stationary data streams, *IEEE Transactions on Neural Networks and Learning Systems* **29**(6): 2516–2529.
- Kaczmarek-Majer, K. and Kiersztyn, A. (2022). Experimental evaluation of the accuracy of an ensemble of fuzzy methods for classification of episodes in bipolar disorder, *IEEE International Conference on Fuzzy Systems (FUZZ-IEEE 2022)*, Padua, Italy, pp. 1–7.
- Kim, Y.S. (2008). Comparison of the decision tree, artificial neural network, and linear regression methods based on the number and types of independent variables and sample size, *Journal of Expert Systems with Application* **34**(2): 1227–1234.
- Klement, E.P., Mesiar, R. and Pap, E. (2000). *Triangular Norms*, Kluwer Academic Publishers, Dordrecht.
- Kowal, M., Skobel, M., Gramacki, A. and Korbicz, J. (2021). Breast cancer nuclei segmentation and classification based on a deep learning approach, *International Journal of Applied Mathematics and Computer Science* **31**(1): 85–106, DOI: 10.34768/amcs-2021-0007.
- Krittanawong, C., Virk, H.U.H., Bangalore, S., Wang, Z., Johnson, K.W., Pinotti, R., Zhang, H., Kaplin, S., Narasimhan, B., Kitai, T., Baber, U., Halperin, J.L. and Tang, W.H.W. (2020). Machine learning prediction in cardiovascular diseases: A meta-analysis, *Scientific Reports* **10**, Article no. 16057.
- Kunapuli, G. (2023). *Ensemble Methods for Machine Learning*, Manning Publications Co, New York.
- Moody, G.B. and Jager, F. (2003). Distinguishing ischemic from non-ischemic ST changes: The PhysioNet/Computers in Cardiology Challenge 2003, *Computers in Cardiology, Thessaloniki, Greece*, pp. 235–237, DOI: 10.1109/CIC.2003.1291134.
- Patro, K.K., Prakash, A.J., Samantray, S., Pławiak, J., Tadeusiewicz, R. and Pławiak, P. (2022). A hybrid

approach of a deep learning technique for real-time ECG beat detection, *International Journal of Applied Mathematics and Computer Science* **32**(3): 455–465, DOI: 10.34768/amcs-2022-0033.

Rak, E., Bazan, J.G., Szczur, A. and Rzaśa, W. (2020). The distributivity law as a tool of k -NN classifiers' aggregation: Mining a cyber-attack data set, *IEEE International Conference on Fuzzy Systems (FUZZ-IEEE 2020), Glasgow, UK*, pp. 1–8.

Rak, E. and Szczur, A. (2021). A comparative assessment of aggregated classification algorithms with the use to mining a cyber-attack data set, *IEEE International Conference on Fuzzy Systems (FUZZ-IEEE 2021), Luxembourg*, pp. 1–6.

Schmidhuber, J. (2015). Deep learning in neural networks: An overview, *Neural Network* **61**: 85–117.

Tanveer, M., Gautam, C. and Suganthan, P.N. (2019). Comprehensive evaluation of twin SVM based classifiers on UCI datasets, *Applied Soft Computing* **83**: 105617.

Zhang, S., Cheng, D., Deng, Z., Zong, M. and Deng, X. (2018). A novel k NN algorithm with data-driven k parameter computation, *Pattern Recognition Letters* **109**: 44–54.

Zipes, D.P., Libby, P., Bonow, R.O., Mann, D.L. and Tomaselli, G.F. (2018). *Braunwald's Heart Disease E-Book: A Textbook of Cardiovascular Medicine*, Elsevier, Amsterdam.



Ewa Rak holds a master's degree in mathematics from the University of Rzeszów (2003) and a PhD degree in mathematics from Pedagogical University in Kraków (2009). She is currently an assistant professor at the University of Rzeszów (Poland), where she is the head of the Department of Functional Equations as well as the Stochastic Processes Laboratory in the Center for Innovation and Transfer of Natural Sciences and Engineering Knowledge. Her main research interests include mathematical modeling, functional equations, aggregation theory, and especially classification methods in machine learning.



Adam Szczur holds a BEng (2012) and an MS (2013) degree in computer science from the University of Rzeszów. He is currently employed in the Institute of Computer Science, University of Rzeszów, Poland. His research interests include rough set theory, data mining, classification methods in machine learning, temporal data, temporal and behavioral patterns.



Jan G. Bazan is an associate professor in the Institute of Computer Science, University of Rzeszów, Poland. He received his MSc degree in mathematics from the Pedagogical University of Rzeszów in 1990, his PhD degree in computer science from Warsaw University in 1999, and his DSc degree in computer science from the Institute of Computer Science, Polish Academy of Sciences, in 2010. His research interests include rough set theory, granular computing, knowledge discovery, data mining, decision-support systems and adaptive systems, especially decision-support systems for medical diagnosis and therapy.



Stanisława Bazan-Socha graduated from Jagiellonian University Medical College in Kraków, Poland, in 1996. She is a specialist in internal medicine, allergy, and clinical immunology, currently employed as an associate professor in the Department of Internal Medicine, Jagiellonian University Medical College in Kraków, Poland. Her research is related to the pathology of asthma and autoimmune diseases as well as artificial intelligence methods in medical data analysis.

Appendix

This appendix presents tables and figures referred to in Section 5.2.

Table A1. Results of $P(KNN)-C_i(KNN)-P_i(KNN)$ and distributivity equation D4.

Distributivity Equation D4				Classifiers' Aggregation Method used $P(kNN)-C_i(kNN)-P_i(kNN)$						
	$\varepsilon=0,1$	$\varepsilon=0,2$	$\varepsilon=0,3$	$\varepsilon=0,4$	$\varepsilon=0,5$	$\varepsilon=0,6$	$\varepsilon=0,7$	$\varepsilon=0,8$	$\varepsilon=0,9$	RAW DATA
TPR(0)	0	0,026	0,105	0,263	0,474	0,763	0,947	1	1	0,842
FPR(0)	0	0	0	0,026	0,105	0,228	0,395	0,675	0,939	0,193
PPV(0)	1	1	1	0,769	0,6	0,527	0,444	0,33	0,262	0,593
TPR(1)	0,237	0,237	0,237	0,237	0,184	0,105	0,105	0,105	0	0,184
FPR(1)	0,272	0,272	0,263	0,211	0,175	0,149	0,07	0,018	0	0,149
PPV(1)	0,225	0,225	0,231	0,273	0,259	0,19	0,333	0,667	0	0,292
TPR(2)	0,711	0,711	0,711	0,658	0,632	0,553	0,447	0,263	0	0,579
FPR(2)	0,465	0,456	0,447	0,447	0,377	0,307	0,211	0,088	0,018	0,298
PPV(2)	0,338	0,342	0,346	0,329	0,358	0,375	0,415	0,5	0	0,393
TPR(3)	0,289	0,289	0,289	0,289	0,289	0,289	0,289	0,211	0,132	0,289
FPR(3)	0,184	0,184	0,175	0,167	0,149	0,079	0,061	0,026	0	0,061
PPV(3)	0,344	0,344	0,355	0,367	0,393	0,55	0,611	0,727	1	0,611
ACC	0,309	0,316	0,336	0,362	0,395	0,428	0,447	0,395	0,283	0,474

Table A2. Results of $P(\text{MLP})-C_i(\text{MLP})-P_i(\text{MLP})$ and distributivity equation D6.

Distributivity Equation D6			Classifiers' Aggregation Method used $P(\text{MLP})-C_i(\text{MLP})-P_i(\text{MLP})$							
	$\epsilon=0,1$	$\epsilon=0,2$	$\epsilon=0,3$	$\epsilon=0,4$	$\epsilon=0,5$	$\epsilon=0,6$	$\epsilon=0,7$	$\epsilon=0,8$	$\epsilon=0,9$	RAW DATA
TPR(0)	0,079	0,105	0,132	0,158	0,553	0,789	0,842	0,895	0,921	0,868
FPR(0)	0	0	0	0	0	0,018	0,018	0,018	0,061	0,018
PPV(0)	1	1	1	1	1	0,938	0,941	0,944	0,833	0,943
TPR(1)	0,763	0,763	0,763	0,763	0,763	0,763	0,763	0,763	0,763	0,763
FPR(1)	0,281	0,281	0,281	0,272	0,167	0,088	0,07	0,061	0,035	0,088
PPV(1)	0,475	0,475	0,475	0,483	0,604	0,744	0,784	0,806	0,879	0,744
TPR(2)	0,895	0,895	0,895	0,895	0,895	0,895	0,895	0,895	0,868	0,868
FPR(2)	0,088	0,079	0,07	0,07	0,053	0,053	0,053	0,053	0,044	0,044
PPV(2)	0,773	0,791	0,81	0,81	0,85	0,85	0,85	0,85	0,868	0,868
TPR(3)	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,921
FPR(3)	0,079	0,079	0,079	0,079	0,07	0,053	0,053	0,044	0,035	0,044
PPV(3)	0,795	0,795	0,795	0,795	0,814	0,854	0,854	0,875	0,897	0,875
ACC	0,664	0,671	0,678	0,684	0,783	0,842	0,855	0,868	0,868	0,855

Table A3. Results of $P(\text{SVM})-C_i(\text{SVM})-P_i(\text{SVM})$ and distributivity equations D4, D6, D8.

Distributivity Equations D4,D6,D8			Classifiers' Aggregation Method used $P(\text{SVM})-C_i(\text{SVM})-P_i(\text{SVM})$							
	$\epsilon=0,1$	$\epsilon=0,2$	$\epsilon=0,3$	$\epsilon=0,4$	$\epsilon=0,5$	$\epsilon=0,6$	$\epsilon=0,7$	$\epsilon=0,8$	$\epsilon=0,9$	RAW DATA
TPR(0)	0,105	0,105	0,105	0,105	0,105	0,816	0,816	0,868	0,895	0,842
FPR(0)	0	0	0	0	0	0,018	0,018	0,018	0,044	0,026
PPV(0)	1	1	1	1	1	0,939	0,939	0,943	0,872	0,914
TPR(1)	0,763	0,763	0,763	0,763	0,763	0,763	0,763	0,763	0,763	0,763
FPR(1)	0,228	0,228	0,228	0,228	0,228	0,061	0,061	0,061	0,053	0,07
PPV(1)	0,527	0,527	0,527	0,527	0,527	0,806	0,806	0,806	0,829	0,784
TPR(2)	0,895	0,895	0,895	0,895	0,895	0,895	0,895	0,895	0,895	0,895
FPR(2)	0,061	0,061	0,061	0,061	0,061	0,044	0,044	0,044	0,044	0,044
PPV(2)	0,829	0,829	0,829	0,829	0,829	0,872	0,872	0,872	0,872	0,872
TPR(3)	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,921
FPR(3)	0,149	0,149	0,149	0,149	0,149	0,079	0,079	0,061	0,035	0,053
PPV(3)	0,673	0,673	0,673	0,673	0,673	0,795	0,795	0,833	0,897	0,854
ACC	0,671	0,671	0,671	0,671	0,671	0,849	0,849	0,862	0,868	0,855

Table A4. Results of $P(\text{NB})-C_i(\text{NB})-P_i(\text{NB})$ and distributivity equation D6.

Distributivity Equation D6			Classifiers' Aggregation Method used $P(\text{NB})-C_i(\text{NB})-P_i(\text{NB})$							
	$\epsilon=0,1$	$\epsilon=0,2$	$\epsilon=0,3$	$\epsilon=0,4$	$\epsilon=0,5$	$\epsilon=0,6$	$\epsilon=0,7$	$\epsilon=0,8$	$\epsilon=0,9$	RAW DATA
TPR(0)	0,105	0,105	0,105	0,105	0,105	0,737	0,737	0,842	0,868	0,737
FPR(0)	0,009	0,009	0,009	0,009	0,009	0,088	0,088	0,105	0,167	0,07
PPV(0)	0,8	0,8	0,8	0,8	0,8	0,737	0,737	0,727	0,635	0,778
TPR(1)	0,5	0,5	0,5	0,5	0,5	0,368	0,368	0,368	0,368	0,658
FPR(1)	0,184	0,184	0,184	0,184	0,184	0,035	0,035	0,035	0,026	0,044
PPV(1)	0,475	0,475	0,475	0,475	0,475	0,778	0,778	0,778	0,824	0,833
TPR(2)	0,684	0,684	0,684	0,684	0,684	0,684	0,684	0,684	0,684	0,868
FPR(2)	0,184	0,184	0,184	0,184	0,184	0,14	0,14	0,114	0,079	0,193
PPV(2)	0,553	0,553	0,553	0,553	0,553	0,619	0,619	0,667	0,743	0,6
TPR(3)	0,842	0,842	0,842	0,842	0,842	0,842	0,842	0,842	0,842	0,658
FPR(3)	0,246	0,246	0,246	0,246	0,246	0,193	0,193	0,167	0,14	0,053
PPV(3)	0,533	0,533	0,533	0,533	0,533	0,593	0,593	0,627	0,667	0,806
ACC	0,533	0,533	0,533	0,533	0,533	0,658	0,658	0,684	0,691	0,73

Table A5. Results of $P(\text{RF})-C_i(\text{RF})-P_i(\text{RF})$ and distributivity equation D5.

Distributivity Equation D5				Classifiers' Aggregation Method used $P(\text{RF})-C_i(\text{RF})-P_i(\text{RF})$						
	$\epsilon=0,1$	$\epsilon=0,2$	$\epsilon=0,3$	$\epsilon=0,4$	$\epsilon=0,5$	$\epsilon=0,6$	$\epsilon=0,7$	$\epsilon=0,8$	$\epsilon=0,9$	RAW DATA
TPR(0)	0,368	0,632	0,842	1	1	1	1	1	1	0,868
FPR(0)	0	0,009	0,044	0,105	0,149	0,193	0,254	0,316	0,482	0,061
PPV(0)	1	0,96	0,865	0,76	0,691	0,633	0,567	0,514	0,409	0,825
TPR(1)	0,763	0,763	0,763	0,684	0,684	0,658	0,605	0,474	0,289	0,763
FPR(1)	0,175	0,132	0,079	0,035	0,009	0	0	0	0	0,07
PPV(1)	0,592	0,659	0,763	0,867	0,963	1	1	1	1	0,784
TPR(2)	0,868	0,868	0,868	0,868	0,842	0,842	0,711	0,658	0,526	0,868
FPR(2)	0,123	0,079	0,053	0,018	0,009	0	0	0	0	0,026
PPV(2)	0,702	0,786	0,846	0,943	0,97	1	1	1	1	0,917
TPR(3)	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,737	0,921
FPR(3)	0,061	0,053	0,026	0,018	0,018	0	0	0	0	0,035
PPV(3)	0,833	0,854	0,921	0,946	0,946	1	1	1	1	0,897
ACC	0,73	0,796	0,849	0,868	0,862	0,855	0,809	0,763	0,638	0,855

Table A6. Results of $P(\text{kNN})-C_i(\text{MLP})-P_i(\text{SVM})$ and distributivity equation D8.

Distributivity Equation D8				Classifiers' Aggregation Method used $P(\text{kNN})-C_i(\text{MLP})-P_i(\text{SVM})$					
	$\epsilon=0,1$	$\epsilon=0,2$	$\epsilon=0,3$	$\epsilon=0,4$	$\epsilon=0,5$	$\epsilon=0,6$	$\epsilon=0,7$	$\epsilon=0,8$	$\epsilon=0,9$
TPR(0)	0,816	0,816	0,816	0,868	0,921	0,947	0,974	0,974	1
FPR(0)	0	0,018	0,026	0,061	0,158	0,193	0,272	0,386	0,649
PPV(0)	1	0,939	0,912	0,825	0,66	0,621	0,544	0,457	0,339
TPR(1)	0,763	0,763	0,763	0,763	0,658	0,579	0,447	0,342	0,263
FPR(1)	0,079	0,079	0,079	0,061	0,044	0,035	0,035	0,026	0,018
PPV(1)	0,763	0,763	0,763	0,806	0,833	0,846	0,81	0,812	0,833
TPR(2)	0,895	0,895	0,895	0,816	0,763	0,763	0,658	0,474	0,184
FPR(2)	0,061	0,053	0,053	0,044	0,044	0,035	0,035	0,026	0,009
PPV(2)	0,829	0,85	0,85	0,861	0,853	0,879	0,862	0,857	0,875
TPR(3)	0,921	0,921	0,921	0,921	0,816	0,816	0,816	0,816	0,5
FPR(3)	0,061	0,053	0,044	0,044	0,035	0,035	0,026	0,026	0,009
PPV(3)	0,833	0,854	0,875	0,875	0,886	0,886	0,912	0,912	0,95
ACC	0,849	0,849	0,849	0,842	0,789	0,776	0,724	0,651	0,487

Table A7. Results of $P(\text{KNN})-C_i(\text{SVM})-P_i(\text{MLP})$ and distributivity equations D1, D8.

Distributivity Equations D1, D8				Classifiers' Aggregation Method used $P(\text{KNN})-C_i(\text{SVM})-P_i(\text{MLP})$					
	$\epsilon=0,1$	$\epsilon=0,2$	$\epsilon=0,3$	$\epsilon=0,4$	$\epsilon=0,5$	$\epsilon=0,6$	$\epsilon=0,7$	$\epsilon=0,8$	$\epsilon=0,9$
TPR(0)	0,789	0,816	0,816	0,842	0,921	0,947	0,974	0,974	1
FPR(0)	0,026	0,035	0,035	0,061	0,158	0,193	0,272	0,377	0,623
PPV(0)	0,909	0,886	0,886	0,821	0,66	0,621	0,544	0,463	0,349
TPR(1)	0,763	0,763	0,763	0,763	0,658	0,579	0,474	0,368	0,289
FPR(1)	0,07	0,07	0,07	0,061	0,035	0,026	0,026	0,026	0,018
PPV(1)	0,784	0,784	0,784	0,806	0,862	0,88	0,857	0,824	0,846
TPR(2)	0,895	0,895	0,895	0,842	0,789	0,789	0,658	0,474	0,211
FPR(2)	0,053	0,044	0,044	0,044	0,044	0,035	0,035	0,026	0,009
PPV(2)	0,85	0,872	0,872	0,865	0,857	0,882	0,862	0,857	0,889
TPR(3)	0,921	0,921	0,921	0,921	0,816	0,816	0,816	0,816	0,5
FPR(3)	0,061	0,053	0,053	0,044	0,035	0,035	0,026	0,026	0,018
PPV(3)	0,833	0,854	0,854	0,875	0,886	0,886	0,912	0,912	0,905
ACC	0,842	0,849	0,849	0,842	0,796	0,783	0,73	0,658	0,5

Table A8. Results of $P(\text{MLP})-C_i(\text{kNN})-P_i(\text{SVM})$ and distributivity equation D5.

Distributivity Equation D5					Classifiers' Aggregation Method used $P(\text{MLP})-C_i(\text{kNN})-P_i(\text{SVM})$				
	$\varepsilon=0,1$	$\varepsilon=0,2$	$\varepsilon=0,3$	$\varepsilon=0,4$	$\varepsilon=0,5$	$\varepsilon=0,6$	$\varepsilon=0,7$	$\varepsilon=0,8$	$\varepsilon=0,9$
TPR(0)	0,868	0,868	0,895	0,895	0,895	0,921	0,947	0,974	1
FPR(0)	0,018	0,018	0,018	0,079	0,123	0,246	0,526	0,833	0,982
PPV(0)	0,943	0,943	0,944	0,791	0,708	0,556	0,375	0,28	0,253
TPR(1)	0,395	0,395	0,395	0,395	0,368	0,289	0,132	0	0
FPR(1)	0,07	0,07	0,061	0,061	0,061	0,061	0,009	0	0
PPV(1)	0,652	0,652	0,682	0,682	0,667	0,611	0,833	0	0
TPR(2)	0,842	0,842	0,842	0,842	0,842	0,842	0,658	0,263	0
FPR(2)	0,298	0,298	0,298	0,289	0,289	0,211	0,123	0,053	0,018
PPV(2)	0,485	0,485	0,485	0,492	0,492	0,571	0,641	0,625	0
TPR(3)	0,5	0,5	0,5	0,395	0,289	0,289	0,211	0,079	0
FPR(3)	0,079	0,079	0,079	0,061	0,061	0,035	0,026	0,009	0
PPV(3)	0,679	0,679	0,679	0,682	0,611	0,733	0,727	0,75	0
ACC	0,651	0,651	0,658	0,632	0,599	0,586	0,487	0,329	0,25

Table A9. Results of $P(\text{MLP})-C_i(\text{SVM})-P_i(\text{kNN})$ and distributivity equations D6, D9, D10 .

Distributivity Equations D6, D9, D10					Classifiers' Aggregation Method used $P(\text{MPL})-C_i(\text{SVM})-P_i(\text{kNN})$				
	$\varepsilon=0,1$	$\varepsilon=0,2$	$\varepsilon=0,3$	$\varepsilon=0,4$	$\varepsilon=0,5$	$\varepsilon=0,6$	$\varepsilon=0,7$	$\varepsilon=0,8$	$\varepsilon=0,9$
TPR(0)	0,026	0,026	0,132	0,132	0,132	0,842	0,895	0,921	1
FPR(0)	0	0	0	0	0	0,009	0,018	0,026	0,36
PPV(0)	1	1	1	1	1	0,97	0,944	0,921	0,481
TPR(1)	0,737	0,737	0,737	0,737	0,737	0,737	0,737	0,737	0,421
FPR(1)	0,237	0,237	0,237	0,237	0,237	0,053	0,044	0,044	0,018
PPV(1)	0,509	0,509	0,509	0,509	0,509	0,824	0,848	0,848	0,889
TPR(2)	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,895	0,605
FPR(2)	0,088	0,088	0,061	0,061	0,061	0,061	0,053	0,053	0,035
PPV(2)	0,778	0,778	0,833	0,833	0,833	0,833	0,854	0,85	0,852
TPR(3)	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,921	0,632
FPR(3)	0,14	0,14	0,132	0,132	0,132	0,07	0,061	0,053	0,035
PPV(3)	0,686	0,686	0,7	0,7	0,7	0,814	0,833	0,854	0,857
ACC	0,651	0,651	0,678	0,678	0,678	0,855	0,868	0,868	0,664

Table A10. Results of $P(\text{SVM})-C_i(\text{kNN})-P_i(\text{MLP})$ and distributivity equation D6.

Distributivity Equation D6					Classifiers' Aggregation Method used $P(\text{SVM})-C_i(\text{kNN})-P_i(\text{MLP})$				
	$\varepsilon=0,1$	$\varepsilon=0,2$	$\varepsilon=0,3$	$\varepsilon=0,4$	$\varepsilon=0,5$	$\varepsilon=0,6$	$\varepsilon=0,7$	$\varepsilon=0,8$	$\varepsilon=0,9$
TPR(0)	0	0	0,447	0,763	0,816	0,842	0,868	0,895	0,921
FPR(0)	0	0	0,009	0,018	0,018	0,018	0,018	0,018	0,368
PPV(0)	0	0	0,944	0,935	0,939	0,941	0,943	0,944	0,455
TPR(1)	0,737	0,737	0,737	0,737	0,737	0,737	0,737	0,737	0,289
FPR(1)	0,211	0,211	0,132	0,079	0,079	0,079	0,079	0,07	0,07
PPV(1)	0,538	0,538	0,651	0,757	0,757	0,757	0,757	0,778	0,579
TPR(2)	0,895	0,895	0,895	0,868	0,868	0,868	0,868	0,868	0,737
FPR(2)	0,298	0,298	0,254	0,211	0,193	0,193	0,193	0,193	0,123
PPV(2)	0,5	0,5	0,54	0,579	0,6	0,6	0,6	0,6	0,667
TPR(3)	0,5	0,5	0,5	0,5	0,5	0,5	0,5	0,5	0,289
FPR(3)	0,114	0,114	0,079	0,07	0,07	0,061	0,053	0,053	0,026
PPV(3)	0,594	0,594	0,679	0,704	0,704	0,731	0,76	0,76	0,786
ACC	0,533	0,533	0,645	0,717	0,73	0,737	0,743	0,75	0,559

Table A11. Results of $P(\text{SVM})-C_i(\text{MLP})-P_i(\text{kNN})$ and distributivity equation D1.

Distributivity Equation	D1				Classifiers' Aggregation Method used $P(\text{SVM})-C_i(\text{MLP})-P_i(\text{kNN})$				
	$\epsilon=0,1$	$\epsilon=0,2$	$\epsilon=0,3$	$\epsilon=0,4$	$\epsilon=0,5$	$\epsilon=0,6$	$\epsilon=0,7$	$\epsilon=0,8$	$\epsilon=0,9$
TPR(0)	0,895	0,895	0,921	0,947	0,974	1	1	1	1
FPR(0)	0,026	0,044	0,167	0,289	0,307	0,465	0,605	0,833	0,939
PPV(0)	0,919	0,872	0,648	0,522	0,514	0,418	0,355	0,286	0,262
TPR(1)	0,737	0,737	0,658	0,5	0,5	0,211	0,132	0,026	0
FPR(1)	0,061	0,061	0,044	0,035	0,026	0,018	0,018	0,018	0,009
PPV(1)	0,8	0,8	0,833	0,826	0,864	0,8	0,714	0,333	0
TPR(2)	0,868	0,816	0,658	0,658	0,658	0,553	0,368	0,105	0
FPR(2)	0,053	0,053	0,053	0,053	0,035	0,026	0,026	0,018	0,009
PPV(2)	0,846	0,838	0,806	0,806	0,862	0,875	0,824	0,667	0
TPR(3)	0,921	0,921	0,842	0,632	0,632	0,632	0,474	0,211	0,132
FPR(3)	0,053	0,053	0,044	0,044	0,044	0,026	0,026	0,018	0
PPV(3)	0,854	0,854	0,865	0,828	0,828	0,889	0,857	0,8	1
ACC	0,855	0,842	0,77	0,684	0,691	0,599	0,493	0,336	0,283

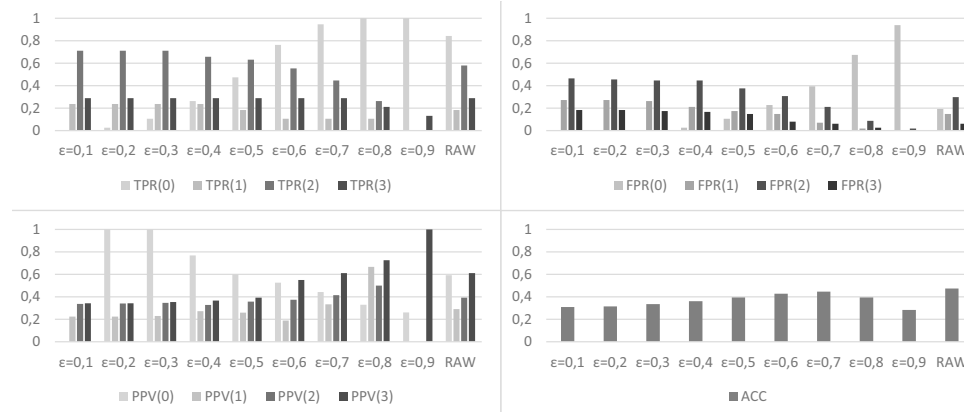


Fig. A1. Results of $P(\text{kNN})-C_i(\text{kNN})-P_i(\text{kNN})$ and distributivity equation D4.

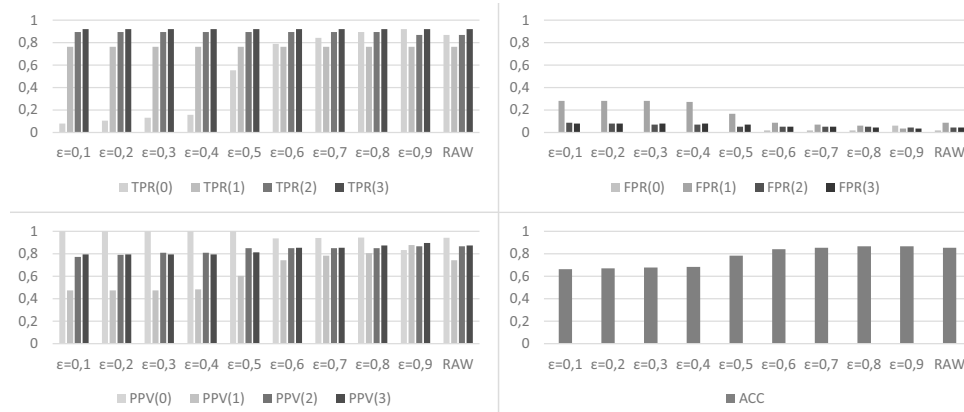


Fig. A2. Results of $P(\text{MLP})-C_i(\text{MLP})-P_i(\text{MLP})$ and distributivity equation D6.

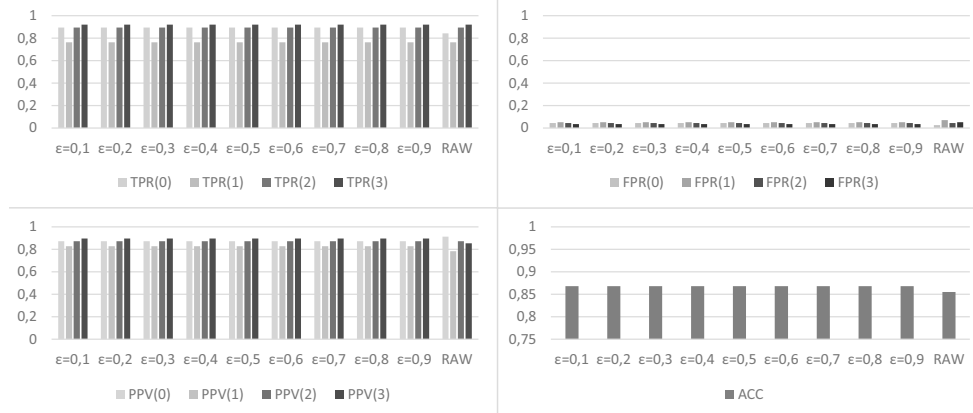


Fig. A3. Results of $P(\text{SVM})-C_i(\text{SVM})-P_i(\text{SVM})$ and distributivity equation D8.

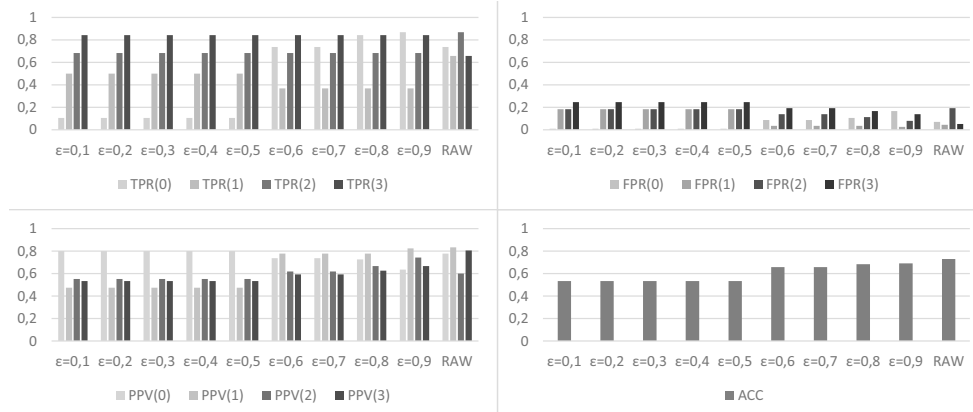


Fig. A4. Results of $P(\text{NB})-C_i(\text{NB})-P_i(\text{NB})$ and distributivity equation D6.

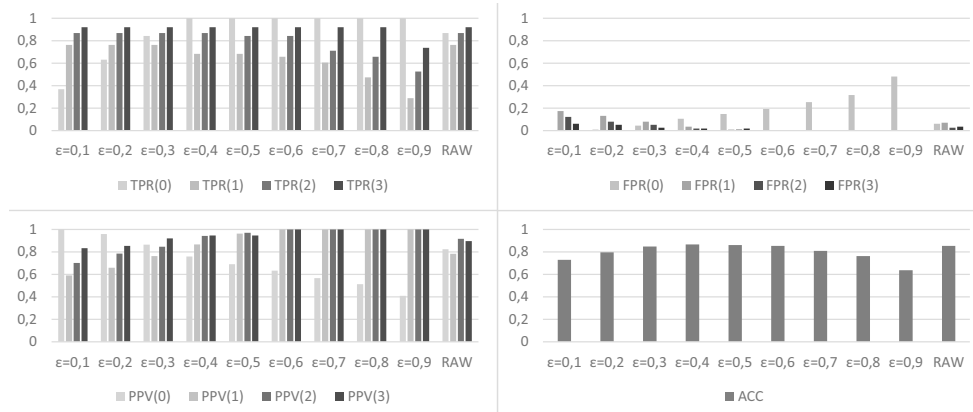


Fig. A5. Results of $P(\text{RF})-C_i(\text{RF})-P_i(\text{RF})$ and distributivity equation D5.

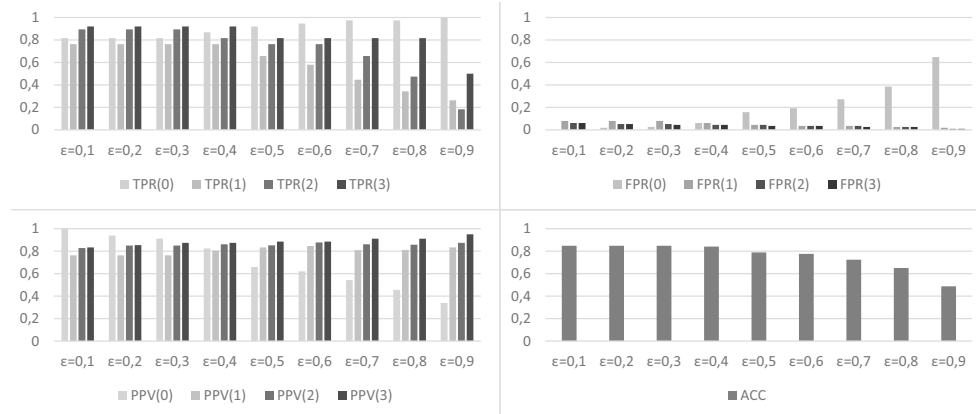


Fig. A6. Results of $P(kNN)-C_i(MLP)-P_i(SVM)$ and distributivity equation D8.

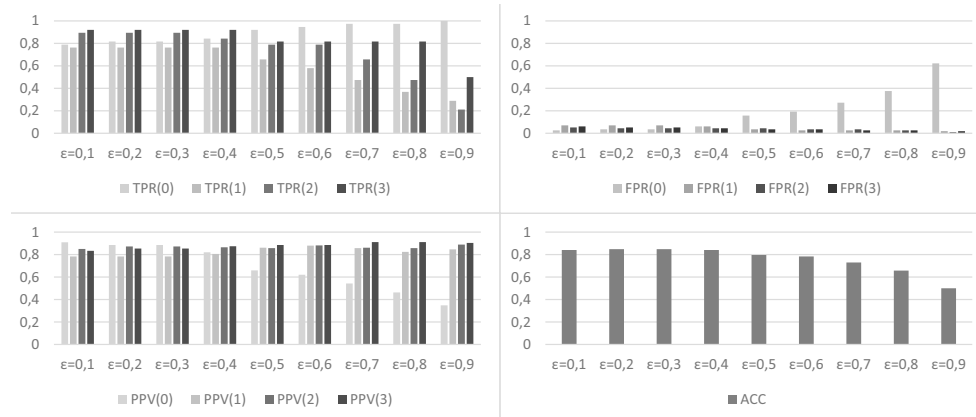


Fig. A7. Results of $P(kNN)-C_i(SVM)-P_i(MLP)$ and distributivity equation D8.

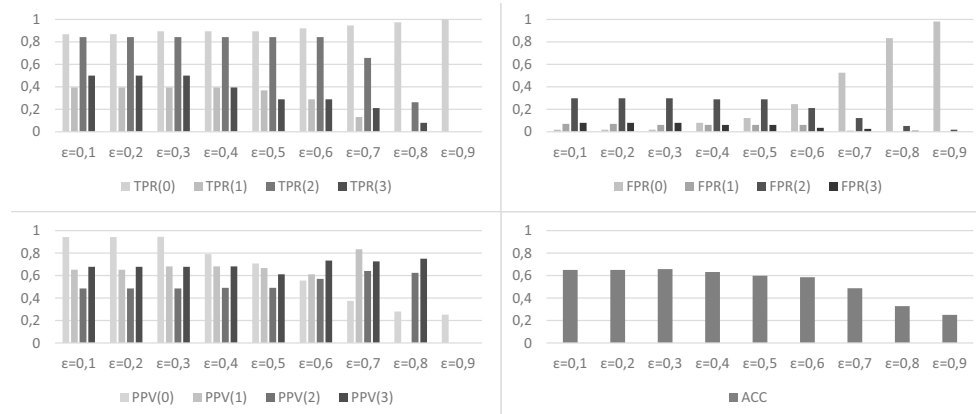
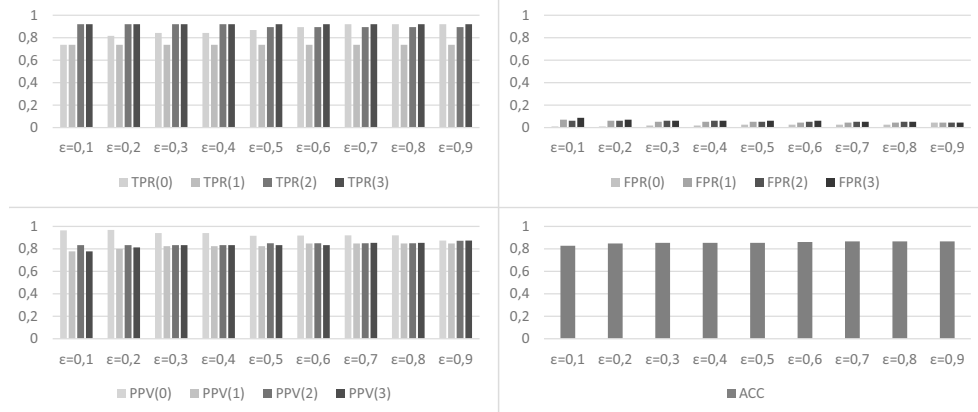
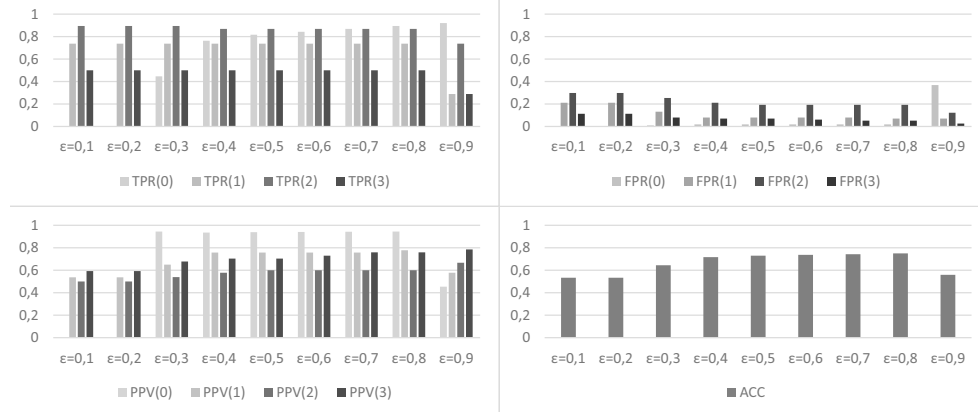
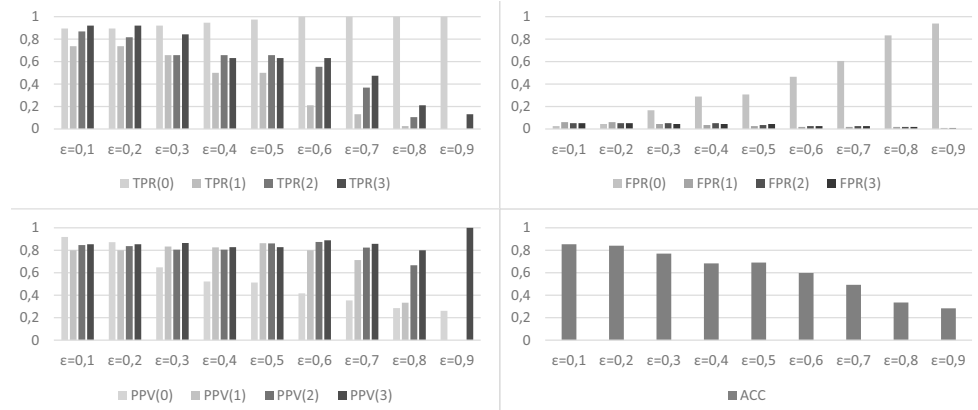


Fig. A8. Results of $P(MLP)-C_i(kNN)-P_i(SVM)$ and distributivity equation D5.

Fig. A9. Results of $P(\text{MLP})-C_i(\text{SVM})-P_i(\text{kNN})$ and distributivity equation D10.Fig. A10. Results of $P(\text{SVM})-C_i(\text{kNN})-P_i(\text{MLP})$ and distributivity equation D6.Fig. A11. Results of $P(\text{SVM})-C_i(\text{MLP})-P_i(\text{kNN})$ and distributivity equation D1.

Received: 17 October 2022

Revised: 15 January 2023

Re-revised: 26 April 2023

Accepted: 13 May 2023