

Unified and Diverse Coalition Formation in Dispersed Data Classification – A Conflict Analysis Approach with Weighted Decision Trees

Małgorzata Przybyła-Kasperek¹[0000–0003–0616–9694] and Jakub Sacewicz¹[0009–0003–1963–7079]

University of Silesia in Katowice, Institute of Computer Science,
Będzińska 39, 41-200 Sosnowiec, Poland
malgorzata.przybyla-kasperek@us.edu.pl, jakub.sacewicz@us.edu.pl

Abstract. The paper delves into the challenge of classification using dispersed data gathered from independent sources. The examined approach involves local models as ensembles of decision trees or random forests constructed based on local data. In the proposed model, a conflict analysis is used to identify the coalitions of local models. Two variants of forming coalitions were checked – unified and diverse – and two different strategies for generating final decisions were explored, allowing one or two of the strongest coalitions to make decisions. The diverse coalition approach is a wholly new and innovative strategy. The methods were tested and compared with corresponding accuracy-based weighted variants. The proposed approach improves classification performance, with weighted variants outperforming unweighted ones in balanced accuracy. Diverse model coalitions are especially effective for challenging and heterogeneous datasets.

Keywords: Dispersed data · Conflict analysis · Decision trees · Random forests · Ensembles of classifiers · Weighted Method.

1 Introduction

In today’s digitalized world, adapting systems to local markets, such as healthcare, banking, and mobile applications, has led to the proliferation of dispersed data. Unlike centralized data systems, where information is gathered in a single repository, dispersed data exists in multiple and diverse environments—from cloud servers to users’ private devices—autonomously gathered without any structural unification. This trend, driven by factors such as law, data security, and the resulting aleatoric uncertainty due to variability and inconsistency across sources, presents both opportunities and challenges that machine learning methods should address, underscoring the growing importance of this research topic.

The topic of dispersed data is most frequently discussed in the distributed learning paradigms [3, 19], which insists on training a group of models separately. Then, various fusion methods, including hierarchical [10] and parallel [2, 15] variants, use the local predictions in final decision-making. However, the process of

making decisions based on ensemble outputs introduces a degree of epistemic uncertainty, especially when the models exhibit conflicting predictions due to insufficient or unrepresentative data in local tables. This uncertainty necessitates robust frameworks to reconcile differences and ensure reliable classification outcomes. The focus is on the diversity among base classifiers [9, 11]. The effectiveness of the global model depends on the method applied to those local classifiers [6, 7]. Mostly, the dependencies between local models are not considered while generating the final prediction. However, that information, in many cases, is crucial and plays a significant role in improving classification quality. Those relations should be taken into account using the selected fusion method. In distributed learning, cooperation, and conflict recognition are rarely applied. Early Artificial Intelligence (AI) conflict studies focused on decision support systems, exploring disputes, identifying key issues, and forming potential coalitions. Various tools have been introduced to analyze conflicts and propose solutions [5]. Nowadays, where there are many multi-agent systems, we can also check their dependencies. Examining them manually, especially when many are dynamic and complex, may be ineffective, burdensome, and unscalable. In this paper, the Pawlak conflict model [12, 13] is considered. The simple model based on rough set theory [14] also gives great insight and understanding of any conflict. The theoretical foundation of rough set theory provides a tool for addressing uncertainty in classification. Rough set theory facilitates the analysis of imprecise or incomplete information by partitioning data into lower and upper approximations, effectively managing uncertainty and conflict in decision-making. The model has been further investigated by many researchers [16, 17, 22], as with some of its developments – the three-way decision theory proposed by Yao [21], with further study by other researchers [20].

This study also uses Pawlak’s model to create a dynamic system in which coalitions of local classifiers are formed. Decision trees and random forests are used as local models, two different variants of creating coalitions are analyzed, and two methods of choosing coalitions (unified and diverse) are used. Additionally, all variants are compared with their weighted variants, where the weights are assigned to each local model based on its accuracy. To the best of our knowledge, the conjunction of forming coalitions with Pawlak’s model with weighted models has never been considered before and compared. The paper shows that, mostly, variants with weights perform better than their corresponding variants. Also, which of the variants of generating and selecting coalitions performs better depends on the chosen data set.

The paper’s contribution is as follows. Introduction of hierarchical decision tree frameworks for dispersed data. Integration of Pawlak’s conflict model to dynamically form coalitions of local models. Exploration of two types of coalitions: unified and diverse. Investigation of two decision strategies: using one or two strongest coalitions. Analysis of the performance of weighted vs. unweighted coalitions. The paper is structured as follows: Section 2 presents the theoretical foundations of the analysis models. Section 3 introduces the methodology, compares results, and discusses findings. The conclusion provides a summary.

2 Methods

Dispersed data is characterized by inconsistencies in structure among different local data sets. One possible approach to address this issue, which also allows data protection, is to deal with each data set in isolation. This is done by generating local models for each local table. In this paper, the ensembles of decision trees or random forests are built. In the following step, we utilize the conflict analysis model to establish coalitions. The process is done for each object dynamically, resulting in different sets of coalitions. Eventually, based on generated coalitions, selected models pass a vote to make the final decision. More formally, we have access to local decision tables represented as $D_i = (U_i, A_i, d)$ for $i \in 1, \dots, n$. In this context, U_i represents the universe, a set of objects; A_i is a set of conditional attributes describing features of the objects; and d denotes the decision attribute, which represents labels. Although some elements may be shared, the objects or attributes may vary. So, the experimental part includes two versions of dispersed data (sets of local tables) – those in which attributes are dispersed and those in which objects are dispersed.

There are various widely known approaches to building local models. This study focuses on a tree-based approach, and we use the decision tree and the random forest models. The models are created using Python with classes from the Skit-learn library (DecisionTreeClassifier or RandomForestClassifier). In future works, we will test other models, as well as the combination of different models, to see if it would result in an improvement in classification quality. For the decision tree, no parameters are specified. For random forest, four different values are selected for the number of estimators: 10, 20, 50, 100. Each local model is trained on a separate data set.

Such ensembles of trees generated based on all local tables make a prediction for the test object $\hat{x}$. The prediction is represented as a vector $[\mu_{i,1}(\hat{x}), \dots, \mu_{i,c}(\hat{x})]$ with dimension c equal to the number of decision classes. The value $\mu_{i,j}(\hat{x})$ is the number of votes cast for decision class j by the random trees in the ensemble. In the paper, we use two approaches – weighted and unweighted methods. During data preprocessing, for weighted approaches, the weights for each local model were calculated with a validation set created randomly with a predefined seed from the test set in a stratified way. In the literature, in most cases, the final decision is made with simple or weighted voting [18] or another fusion method [6]. The novelty of this work is using prediction vectors in conjunction with the conflict analysis method. This method is used to find coalitions that influence the final decision. Two types of local models' coalitions are considered in this study: coalitions of local models with similar opinions and coalitions of local models with diverse opinions are considered. Both approaches are relevant and based on different justifications. The first approach seeks consensus opinions, assuming that the majority of local models are accurate and make correct decisions. The second emphasizes diversity, drawing on the ensemble of classifiers principle, which suggests that varied local models enhance classification quality. For the task of creating those coalitions, Pawlak's conflict model [12, 13] is applied.

In Pawlak's conflict analysis model, information about the conflict situation is stored in an information system $S = (LM, V)$, where LM is a set of local models, and V is a set of decision attribute values. Function $v : LM \rightarrow \{-1, 0, 1\}$ for each $v \in V$ and $i \in LM$ is defined

$$v(i) = \begin{cases} 1 & \text{if the coordinate } \mu_{i,v}(\hat{x}) \text{ for decision } v \text{ in the prediction vector of} \\ & \text{model } i \text{ has the maximum value of all coordinates in this vector.} \\ 0 & \text{if the coordinate } \mu_{i,v}(\hat{x}) \text{ for decision } v \text{ in the prediction vector of} \\ & \text{model } i \text{ is this vector's second highest of all coordinates.} \\ -1 & \text{in other cases} \end{cases} \quad (1)$$

That is a way of storing opinions of local models on the test object. For each model, the decision with the most significant support is assigned the value 1, which is interpreted as supporting the decision by the model. The next most supported decision is assigned the value 0, which means the model is neutral to this decision. We assign the value -1 for all other decisions, which means the model is against them. A conflict function is applied in the next step of the Pawlak conflict analysis model. The conflict function $\rho : LM \times LM \rightarrow [0, 1]$ is defined as follows:

$$\rho(i, j) = \frac{\text{card}\{v \in V : v(i) \neq v(j)\}}{\text{card}\{V\}}, \quad (2)$$

where $\text{card}\{V\}$ is the cardinality of the set of decision classes and $i, j \in LM$.

For unified coalitions, set $X \subseteq LM$ is a coalition of local models if for every $i, j \in X$ we have $\rho(i, j) < 0.5$. The coalition will include those local models for which the opinion is consistent in more than half of the decision classes. The conditions are opposite in diverse coalitions; they include local models inconsistent with more than half of the decision classes. Coalitions do not have to be disjoint sets. One local model can belong to many coalitions simultaneously, reflecting real arrangements in everyday life. Algorithm 1 presents a pseudo-code for determining coalitions.

Once the coalitions are generated, the next step is to make a final decision using them. This work explores four different approaches combined with both ways of creating coalitions described above.

The first approach assumes using the **two strongest coalitions** (in terms of the number of models in the coalition). The chosen coalitions are the ones with the maximum number of members. In situations where there are coalitions with the size of the second strongest coalition, those coalitions are also considered. Then, the prediction vectors of all chosen models are added, and a single prediction vector is formed. If some models are part of many coalitions, their vectors are included every time they occur. Finally, the decision with the biggest score is made – the decision with the largest value of the coefficient of the joint prediction vector. The second approach is analogous to the previous one, except that we choose only **one most strongest coalition**, which will decide.

Both **weighted** approaches work analogously to unweighted ones. Besides that, each local model's prediction vector is multiplied by the weight, equal to

Algorithm 1: Create Coalitions

```

Data:  $S = (LM, V)$ 
Result:  $Coalition\_set$ 
1 begin
2    $Coalition\_set \leftarrow$  empty list;           // Initialize an empty list
3    $Boolean \leftarrow TRUE$ ;
4   foreach  $X \subseteq LM$  do
5     foreach  $i, j \in X$  do
6       if  $\rho(i, j) \geq 0.5$  then
7         // for unified coalitions or  $\rho(i, j) \leq 0.5$  for diverse coalitions
8          $Boolean = FALSE$ 
9       if  $Boolean = TRUE$  then
10         $Coalition\_set \leftarrow X$ 
11 return  $Coalition\_set$ 
    
```

the model's accuracy, estimated by evaluating the model using the validation set. Before testing the global model, the test set is stratified into test and validation sets with proportions of 0.5 to 0.5.

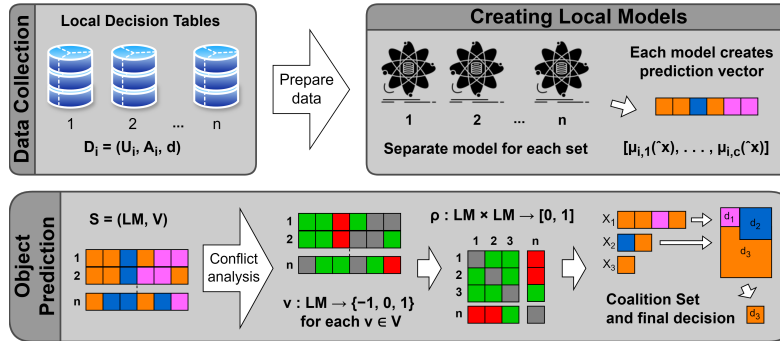


Fig. 1. Model generation stages.

Figure 1 shows the proposed hierarchical framework with conflict analysis for ensembles of local models (decision trees of random forests), presenting stages discussed above. At the beginning, the dispersed data are given. For each local table, we build a separate predictive model. During the classification stage, we retrieve the prediction vector from every local model. Those vectors are next used to create an information system representing the conflict situation. Then, the coalitions are formed, and depending on the method, one or some of them are selected to make the final decision. The last step is to calculate the score for each decision class and choose the one with the highest score.

3 Data sets and experimental results

The study evaluated approaches using three datasets from the UC Irvine Machine Learning Repository [1, 23, 8] and one empirical dataset [4]. The Avian Influenza dataset tracks global human infections from 12 countries, organized into four tables by region (e.g., Egypt, Vietnam, Indonesia) with region-specific attributes. Its structure supports predictive modeling and epidemiological research to understand avian influenza’s impact on health. The Car Evaluation dataset includes six categorical features (e.g., price, maintenance, doors, capacity, luggage size, safety) and classifies cars into four categories, making it suitable for machine learning evaluation. The Lymphography dataset classifies lymphatic diseases using features like node appearance and histological findings to aid medical diagnosis. The Vehicle Silhouettes dataset distinguishes vehicle types based on shape features from silhouettes. Dataset characteristics are summarized in Table 1.

Table 1. Data set characteristics

Data set	# The training set	# The test set	# Conditional attributes	Attributes type	# Decision classes	Source
Avian influenza	205	89	5	Categorical and Integer	4	[4]
Car evaluation	1210	518	6	Categorical	4	[1]
Lymphography	104	44	18	Categorical	4	[23]
Vehicle Silhouettes	592	254	18	Integer	4	[8]

The research investigates two different methodologies concerning data dispersion with respect to both attribute and object dispersion. All datasets within the UCI repository were consolidated into one table, which was subsequently dispersed across various local tables (3, 5, 7, 9, and 11 local tables for each dataset). For the Lymphography and the Vehicle Silhouettes, attributes were dispersed randomly across the tables. Efforts were made to balance the number of attributes allocated to each table. The same objects are present in all tables, but their identifiers were omitted to simulate a real-world scenario where recognizing identical objects across tables is not possible. The Car dataset was dispersed relative to the objects. This means that the objects were divided among the tables in a stratified and random manner, and all attributes were included in each table. The Avian dataset was collected in 12 countries and divided into four local tables based on what country the object comes from. Local tables include countries like Egypt, Vietnam, Indonesia, and others were created. Each table contains conditional attributes and objects from specific countries.

The evaluation insists on repeating a test 10 times to check the stability of all methods. Despite the random behavior of some approaches, to make the results reproducible, the set of 10 seeds was randomly chosen, and each test was conducted with its assigned seed, starting from learning local models on local tables, generating validation sets, and ending on generating global decisions by all approaches.

The assessment of classification quality relied on the test set with various accuracy measures: Classification Accuracy (Acc), Accuracy’s standard devia-

tion (Acc SD), Recall, Precision (Prec.), balanced accuracy (BAcc), balanced accuracy standard deviation (BAcc SD), and F1 measure (F1). This study uses two methods to calculate the F1 measure: weighted and macro.

The results are summarized in Figures 2–5, presenting averaged metric values from all 10 achieved metrics with standard deviation for accuracy and balanced accuracy. Each dataset, characterized by varying degrees of dispersion, is evaluated using three approaches: without coalitions, with coalitions representing agents' agreement, and with coalitions representing agents' disagreement. The following notations are used:

- Probability Sum; normal – Local classifiers (random forest or decision tree) independently predict; their prediction vectors are summed, and the decision with the highest value is selected;
- Probability Sum; weighted – Similar to normal, but classifiers are weighted by accuracy estimated on a validation set (split from the test set);
- Unified groups; one strongest – Coalitions are formed based on prediction vectors using classical Pawlak's approach. The strongest coalition's summed vectors determine the decision;
- Unified groups; two strongest – Extends the above by considering the two strongest coalitions;
- Unified groups; weighted one strongest/weighted two strongest – Combines the above with accuracy-based weighting of coalition vectors;
- Diverse groups; one strongest/two strongest/weighted one strongest/weighted two strongest – Forms coalitions of local models with conflicting predictions (vector distance greater than 0.5). Variants include one/two strongest or weighted coalitions.

Additionally, various numbers of estimators were tested for the random forest classifier, specifically 10, 20, 50, and 100. The tables present the best result achieved, along with the corresponding number of estimators that produced this outcome. For each dataset, the table indicates the best result obtained.

Based on the results in Figures 2–5, we can draw the following conclusions. The Avian Influenza dataset showed its highest performance using the random forest (RF) classifier with the 'Diverse groups' approach. This result suggests that conflicting classifiers through coalition formation improve predictive accuracy. We believe that this is due to the variety of approaches taken, as differing perspectives significantly impact quality. The ability to aggregate diverse predictions likely mitigates inconsistencies caused by object dispersion across multiple tables (countries). The decision tree (DT) classifier using the 'Unified groups' or 'Diverse groups' approach consistently gives the best results for most versions of dispersion. The Car Evaluation dataset has categorical features, making it well-suited for decision trees. Interestingly, for the Car dataset, the 'Probability Sum' approach, without coalitions, also often produces good results (for dispersion 3LT, 7LT, 9LT). It is also worth noting that, for this dataset, the use of weights does not lead to any improvements. The Lymphography dataset, characterized by complex medical features, performed best when using weighted random forests in the 'Diverse groups' configuration. This suggests that classifier

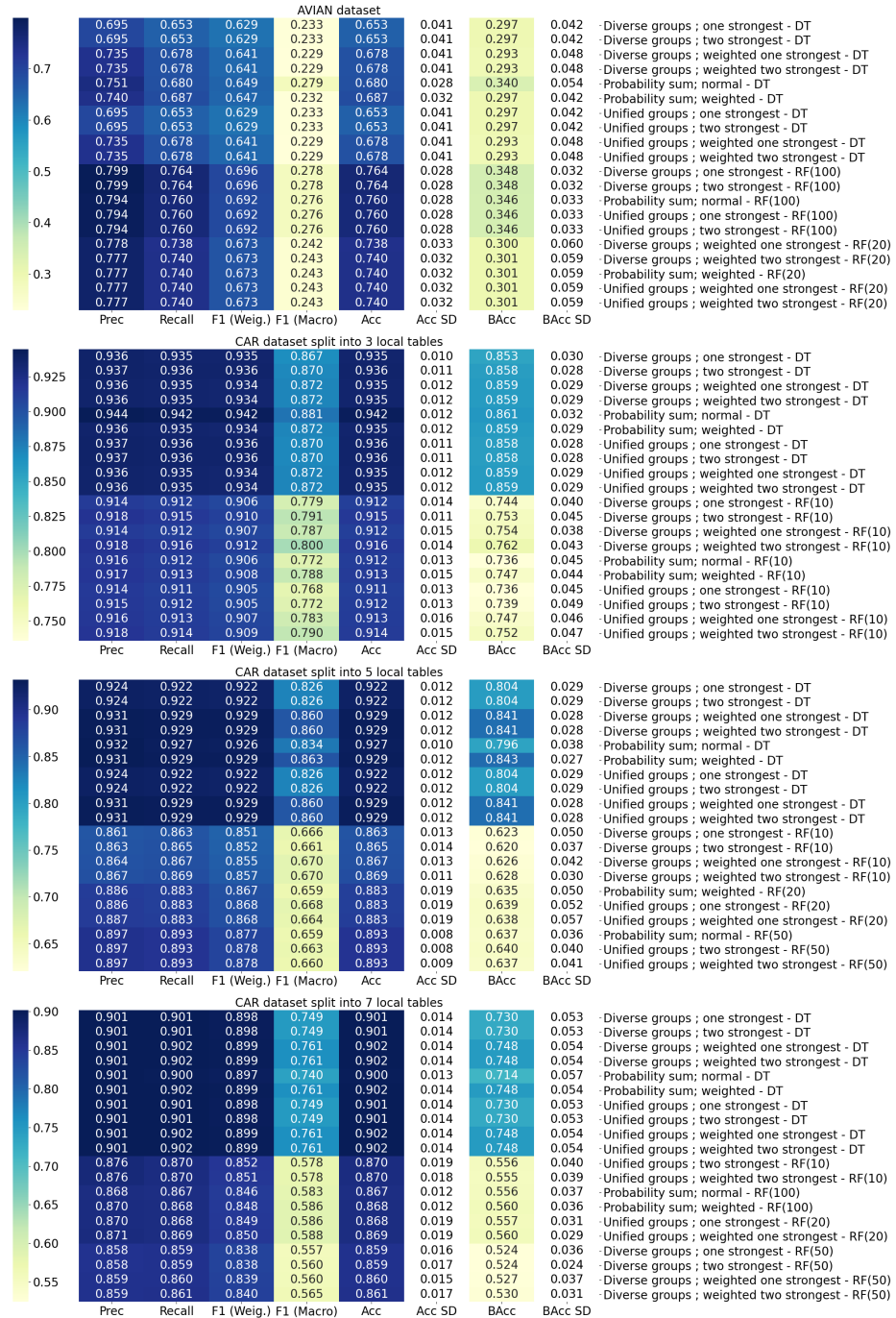


Fig. 2. Results of precision (Prec.), recall, F-measure (F-m.), balanced accuracy (*bacc*) and classification accuracy (*acc*) for the considered approaches Part 1. RF is the abbreviation Random Forest and DT for Decision Tree.

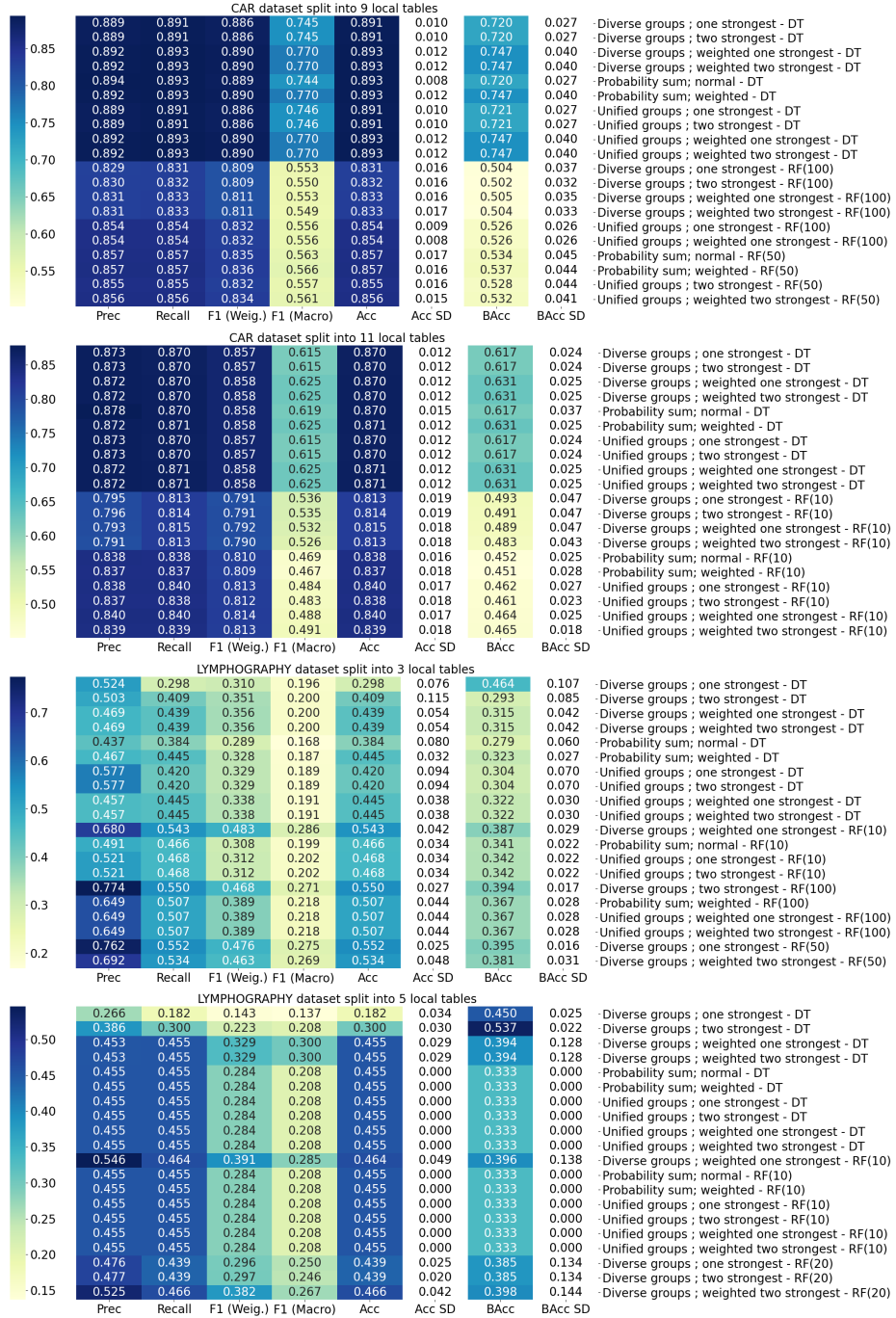


Fig. 3. Results of precision (Prec.), recall, F-measure (F-m.), balanced accuracy (*bacc*) and classification accuracy (*acc*) for the considered approaches Part 2. RF is the abbreviation Random Forest and DT for Decision Tree.

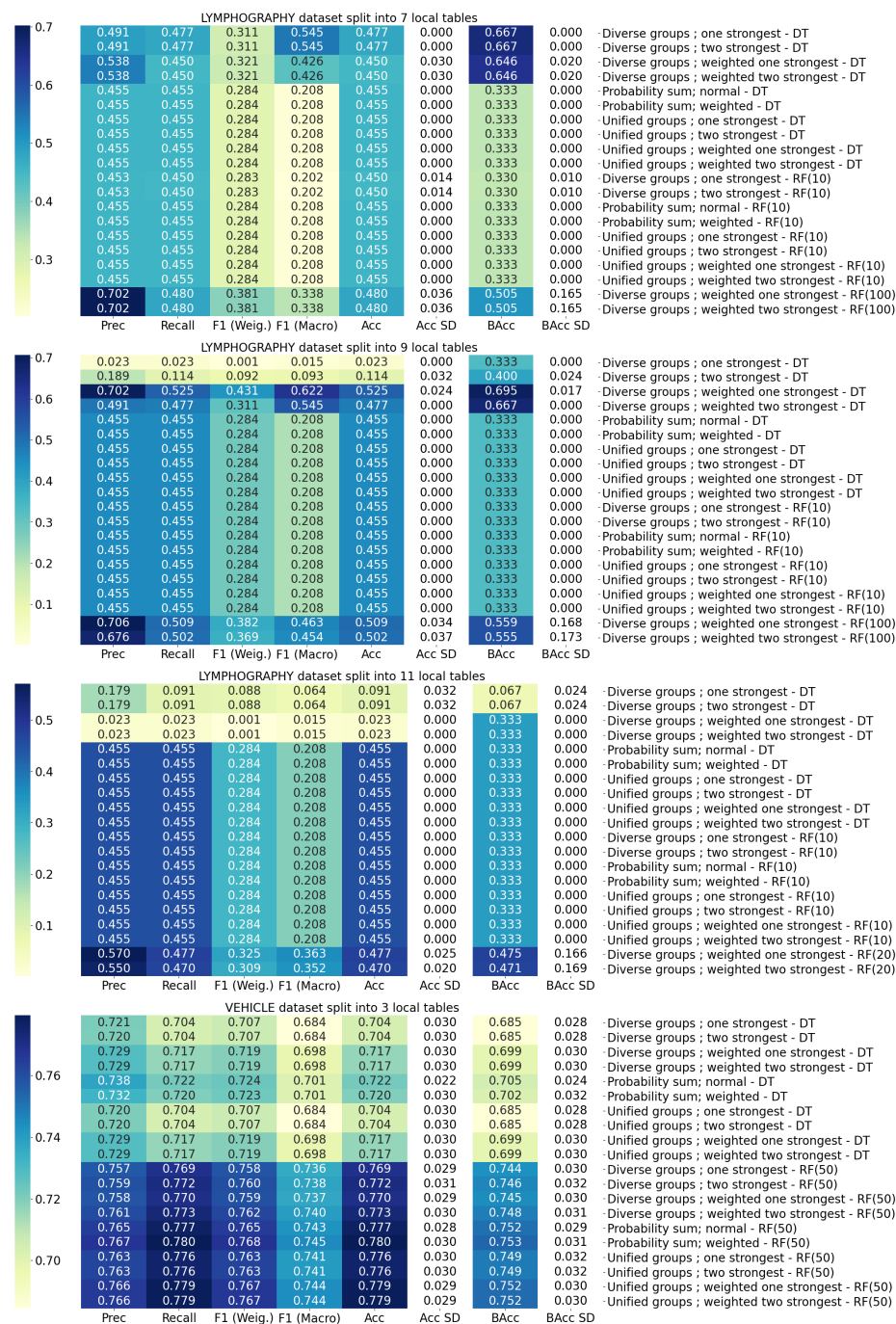


Fig. 4. Results of precision (Prec.), recall, F-measure (F-m.), balanced accuracy (*bacc*) and classification accuracy (*acc*) for the considered approaches [Part 3](#). RF is the abbreviation Random Forest and DT for Decision Tree.

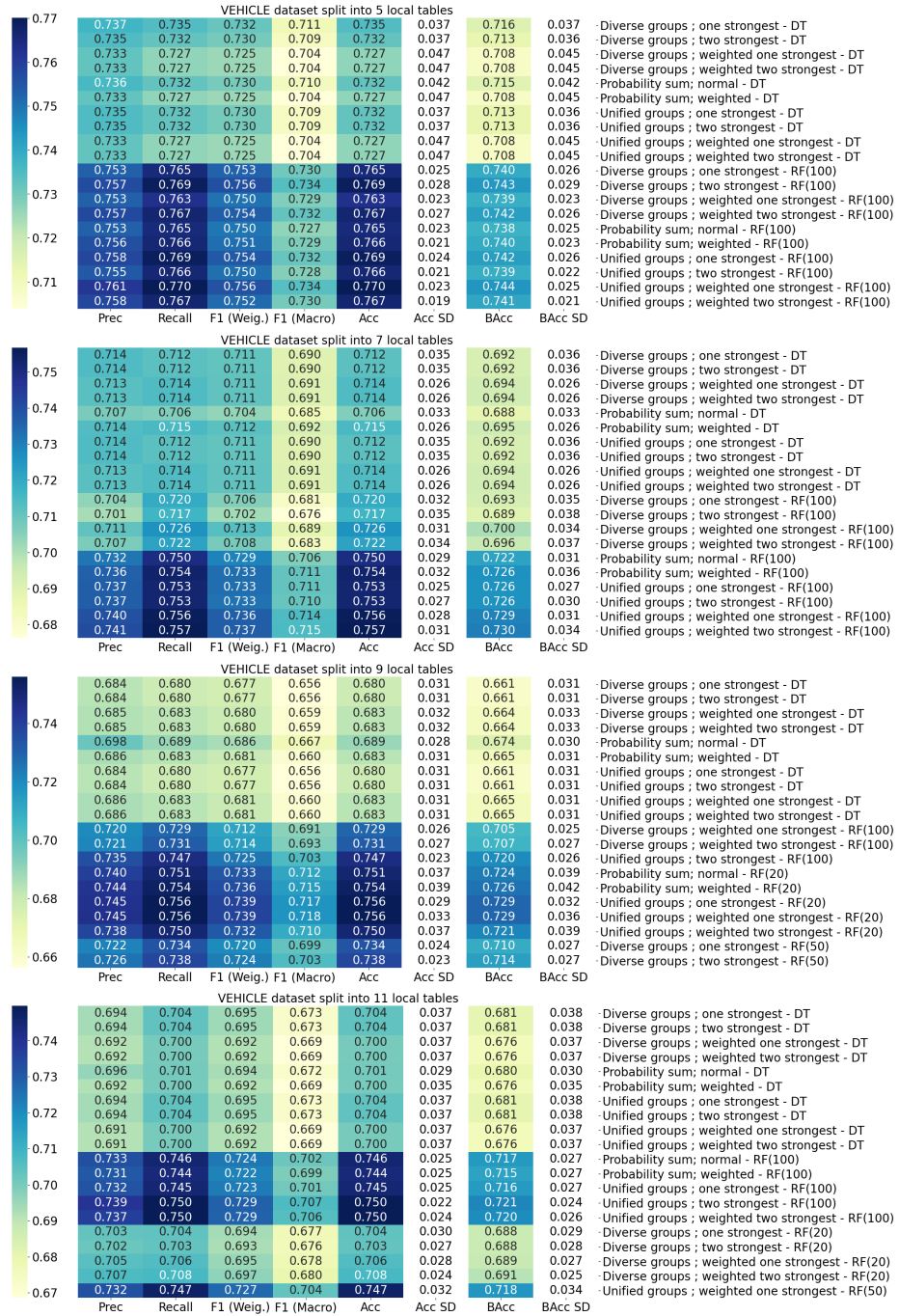


Fig. 5. Results of precision (Prec.), recall, F-measure (F-m.), balanced accuracy (*bacc*) and classification accuracy (*acc*) for the considered approaches Part 4. RF is the abbreviation Random Forest and DT for Decision Tree.

weighting based on validation set accuracy compensates for prediction conflicts. However, results were generally lower due to the dataset’s challenging nature. Given that decision trees sometimes obtain the highest prediction quality (especially for 9LT), it is important to emphasize that accurate weighting and conflict resolution are crucial for this dataset. For the Vehicle Silhouettes dataset, the best-performing approach was random forests with weighted and sometimes unweighted ‘Unified groups’.

In general, random forests consistently outperformed decision trees in datasets with numerical or complex features, such as Vehicle Silhouettes and Lymphography. In contrast, decision trees excelled in the Car Evaluation dataset due to its well-defined categorical attributes. It could also be a result of a number of features selected to build trees. By default, a single decision tree uses all features, while in a random forest, the square root of the number of features. Higher dispersion levels negatively impacted all approaches, especially in complex datasets like Lymphography. Approaches utilizing weighted classifiers or diverse coalition strategies mitigated this decline. Forming coalitions among classifiers with conflicting predictions improved results, emphasizing the value of ensemble diversity. Applying classifier accuracy-based weights led to better results in most cases.

Statistical tests were performed to confirm the observed validities and balanced accuracy values were used for comparison. At first, the balanced accuracy values of all twenty approaches were compared – so twenty dependent samples each containing of 16 observations were created, representing the results for each dataset and dispersion version. Since the balanced accuracy is ratio-scaled and normal distribution is not confirmed, also the samples are small the Friedman test was used to determine whether the differences in balanced accuracy values among the approaches were statistically significant. The Friedman test indicated that there is no statistically significant difference in mean balanced accuracy among the twenty approaches, $\chi^2(15, 19) = 9.58, p = 0.96$. A comparative box plot illustrating the balanced accuracy results for the twenty methods is provided in Figure 6. Although the statistical test did not confirm the significance of the mean differences for such a large number of samples, the graph clearly shows that the results for the weighted approaches are significantly higher than those for the other methods. Therefore, further statistical tests were conducted.

Next, we analyze the differences in average balanced accuracy among the four approaches: unified groups unweighted and weighted, diverse groups unweighted and weighted. This time, the results were organized into 4 dependent groups, each containing 64 observations. Similarly, the Friedman test was conducted on balanced accuracy values. The test confirmed a statistically significant difference in the averages among at least two of the approaches, $\chi^2(64, 3) = 17.75, p = 0.0005$. To pinpoint the specific differences, a post-hoc Dunn-Bonferroni test was performed, with the significant results highlighted in blue in Table 2. The test revealed significant differences between the unweighted and weighted approaches. A comparative box plot (Figure 7) shows that the balanced accuracy results are slightly better for weighted than unweighted approaches. This was also confirmed by the Wilcoxon test (p-value 0.0001) for results organized into two groups –

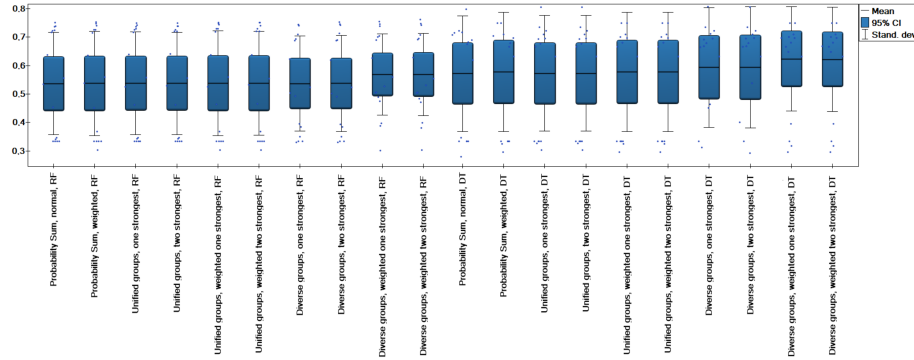


Fig. 6. Comparison of balanced accuracy obtained for all analyzed approaches.

unweighted and weighted approaches – with 128 observations in each. It was checked analogously with the Wilcoxon test (p-value 0.17) that the division into approaches with different methodologies for forming coalitions – unified groups and diverse groups – does not bring significant differences. However, when we analyzed individual data sets, it was apparent that sometimes the approach with consensus coalitions is better, while for difficult data sets, the approach with incompatible coalitions is better. Thus, the effectiveness of the coalition formation approach depends on the dataset. However, in general, it can be concluded that weighted approaches yield better results.

Table 2. p-values for the post-hoc Dunn Bonferroni test for approaches: unified groups unweighted and weighted; diverse groups unweighted and weighted

p-value	Unified groups	Unified groups weighted	Diverse groups	Diverse groups weighted
Unified groups		0.04	1	0.04
Unified groups, weighted	0.04		0.03	1
Diverse groups	1	0.03		0.03
Diverse groups, weighted	0.04	1	0.03	

In conclusion, the results indicate the importance of classifier accuracy-based weighting for dispersed data. Diverse coalition strategies, which group classifiers with conflicting predictions, proved particularly effective for datasets with complex features or high dispersion, such as Lymphography and Vehicle Silhouettes. In contrast, unified coalition approaches often performed better in datasets with categorical features, exemplified by the Car Evaluation dataset. Overall, the findings emphasize the critical role of classifier diversity and weighting in achieving robust predictive performance across varied datasets.

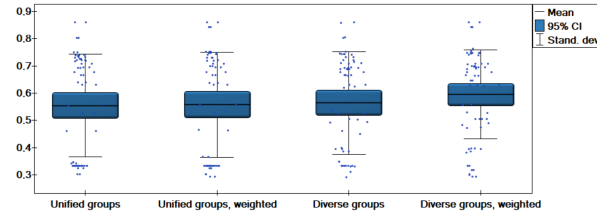


Fig. 7. Comparison of balanced accuracy obtained for approaches: unified groups unweighted and weighted; diverse groups unweighted and weighted.

4 Conclusion

This paper proposes a hierarchical classification model based on dispersed data. Decision trees and random forests were used with conflict model analysis. For the first time, a model with diverse coalitions was used in conjunction with the sum method. In addition to that, the weighted variants were introduced and compared with the unweighted ones. In this work, tests were made on sixteen different dispersed datasets, some with respect to objects and others with respect to attributes. The results were compared with other known literature methods: the sum and weighted sum methods with trees or random forest methods.

The proposed approach yields better results for most of the tested data sets. It was also statistically proven that in terms of classification quality, determined by balanced accuracy measure, weighted variants provide better classification quality than corresponding variants without assigning weights for local models. Also, which method of creating coalitions is better depends on the data set — coalitions of diverse models enhance classification quality for more challenging and diverse data sets.

In future work, the k-nearest neighbors, AdaBoost classifiers, and multilayer perceptrons are planned to be used as local classifiers in conjunction with the conflict analysis and coalition formation method and with other methods of determining differences among local models. The plans also includes features importance and their similarity between coalitions.

References

1. Bohanec, M.: (1997). Car Evaluation. UCI Machine Learning Repository. <https://doi.org/10.24432/C5JP48>.
2. Czarnowski, I. Weighted Ensemble with one-class Classification and Over-sampling and Instance selection (WECOI): An approach for learning from imbalanced data streams, *Journal of Computational Science*, 61, (2022) 101614, ISSN 1877–7503.
3. Czarnowski, I., Jędrzejowicz, P. Ensemble online classifier based on the one-class base classifiers for mining data streams. *Cybernetics and Systems*, **2015**, 46(1-2), 51–68.
4. Fiebig, L., Soyka, J., Buda, S., Buchholz, U., Dehnert, M., Haas, W., 2011, Avian influenza A(H5N1) in humans - line list, <http://dx.doi.org/10.25646/7661> (accessed on 15 February 2024).

5. Giordano, R., Passarella, G., Uricchio, V. F., Vurro, M. (2007). Integrating conflict analysis and consensus reaching in a decision support system for water resource management. *Journal of environmental management*, 84(2), 213-228.
6. Kashinath, S. A.; Mostafa, S. A.; Mustapha, A.; Mahdin, H.; Lim, D.; Mahmoud, M. A.; Mohammed, M.A.; Al-Rimy, B.A.S.; Fudzee M. F.; Yang, T. J. Review of data fusion methods for real-time and multi-sensor traffic flow analysis. *IEEE Access*, (2021) 9, 51258-51276.
7. Kuncheva, L. I. Combining pattern classifiers: methods and algorithms. 2014. John Wiley & Sons.
8. Mowforth, P. Shepherd, B. Statlog (Vehicle Silhouettes) [Dataset]. UCI Machine Learning Repository. <https://doi.org/10.24432/C5HG6N>.
9. Nam, G.; Yoon, J.; Lee, Y.; Lee, J. Diversity matters when learning from ensembles. *Advances in Neural Information Processing Systems*, (2021) 34, 8367-8377.
10. Ng, W. W.; Zhang, J.; Lai, C. S.; Pedrycz, W.; Lai, L. L.; Wang, X. Cost-sensitive weighting and imbalance-reversed bagging for streaming imbalanced and concept drifting in electricity pricing classification. *IEEE Transactions on Industrial Informatics*, 15(3), 1588-1597, (2018).
11. Ortega, L. A.; Cabañas, R.; Masegosa, A. Diversity and Generalization in Neural Network Ensembles. In *International Conference on Artificial Intelligence and Statistics* (2022) (11720-11743). PMLR.
12. Pawlak, Z. Some remarks on conflict analysis. *Eur. J. Oper. Res.* **2005**, 166, 649-654.
13. Pawlak, Z. Conflict analysis. In *Proceedings of the Fifth European Congress on Intelligent Techniques and Soft Computing (EUFIT'97)*, Aachen, Germany, 8-12 September 1997; pp. 1589-1591.
14. Pawlak, Z. (1982). Rough sets. *International journal of computer & information sciences*, 11, 341-356.
15. Pławiak, P.; Abdar, M.; Pławiak, J.; Makarenkov, V.; Acharya, U. R. DGHNL: A new deep genetic hierarchical network of learners for prediction of credit scoring. *Information Sciences*, 516, 401-418, (2020).
16. Przybyła-Kasperek, M.; Sacewicz, J. (2024). Ensembles of random trees with coalitions-a classification model for dispersed data. *Procedia Computer Science*, 246, 1599-1608.
17. Przybyła-Kasperek, M.; Wakulicz-Deja, A. (2014). A dispersed decision-making system-The use of negotiations during the dynamic generation of a system's structure. *Information Sciences*, 288, 194-219.
18. Tasci, E., Uluturk, C., Ugur, A.: A voting-based ensemble deep learning method focusing on image augmentation and preprocessing variations for tuberculosis detection. *Neural Computing and Applications*, 33(22), 15541-15555, (2021).
19. Verbeaeken, J.; Wolting, M.; Katzy, J.; Kloppenburg, J.; Verbelen, T.; Rellermeyer, J. S. A survey on distributed machine learning. *ACM computing surveys*, 53(2), 1-33, (2020).
20. Xiaonan Li, Yucong Yan. (2024). A dynamic three-way conflict analysis model with adaptive thresholds, *Information Sciences*, Volume 657.
21. Yao, Y. (2010). Three-way decisions with probabilistic rough sets. *Information sciences*, 180(3), 341-353.
22. Zohaib Gillani, Zia Bashir, Saira Aquil. A game theoretic conflict analysis model with linguistic assessments and two levels of game play. *Information Sciences*, 677, (2024) 120840, ISSN 0020-0255.
23. Zwitter, M. and Soklic, M.: (1988). Lymphography. UCI Machine Learning Repository. <https://doi.org/10.24432/C54598>.