



Research Article

Comparing ECLAT and Decision Tree for Drug Therapy Recommendation Rules on Multi Label Clinical Data

Muh Rayhan Fahreza^{1*}; Harlinda²; Herdianti Darwis³; Roesman Ridwan Raja⁴

¹ Universitas Muslim Indonesia, Makassar, Indonesia, rayhanfahreza56@gmail.com

² Universitas Muslim Indonesia, Makassar, Indonesia, harlinda@umi.ac.id

³ Universitas Muslim Indonesia, Makassar, Indonesia, herdianti.darwis@umi.ac.id

⁴ Kyushu Institute of Technology, Iizuka, Japan, raja.roesman-ridwan757@mail.kyutech.jp

Correspondence should be addressed to Muh Rayhan Fahreza; rayhanfahreza56@gmail.com

Received 29 March 2026; Accepted 30 April 2026; Published 31 July 2026

© Authors 2026. CC BY-NC 4.0 (non-commercial use with attribution, indicate changes).

License: <https://creativecommons.org/licenses/by-nc/4.0/> — Published by Indonesian Journal of Data and Science.

Abstract:

Introduction: Accurate sorting of plastic waste using Resin Identification Codes (RICs) is essential for improving recycling quality. However, conventional deep learning approaches generally require large labeled datasets, which are difficult and costly to collect for small RIC symbols on plastic packaging. **Method:** This study employed a Few-Shot Learning approach based on Prototypical Networks using a 7-way 5-shot episodic training configuration. A self-collected dataset of 350 images covering seven RIC categories was used, and three backbone architectures, ConvNet4, ResNet-18, and EfficientNet-B2, were compared. An ablation study evaluated support-set augmentation, followed by supervised fine-tuning of the selected model. **Results and Discussion:** EfficientNet-B2 achieved the highest episodic accuracy of 93.36%, outperforming ResNet-18 at 88.71% and ConvNet4 at 54.14%. EfficientNet-B2 with light augmentation achieved 85.71% accuracy on the fixed 42-image test set. Most errors occurred between visually similar HDPE and PP symbols. Fine-tuning corrected five of six misclassifications, increasing test accuracy to 90.48% and the F1-score from 0.857 to 0.903. **Conclusion:** Prototypical Networks with an EfficientNet-B2 backbone and cosine distance provide an effective approach for RIC classification under limited-data conditions and offer a practical foundation for automated plastic-waste sorting systems.

Keywords: ECLAT; Decision Tree; Association Rule Mining; Drug Recommendation; Clinical Decision Support

Dataset link: [BPJS Drug Recommendation Dataset at the Bonto Perak Community Health Center, Pangkep Regency, South Sulawesi.](#)

1. Introduction

The digitalization of medical records has resulted in a vast volume of clinical data; however, leveraging this data to support objective drug therapy decisions remains a challenge. Selecting the appropriate drug therapy is crucial, as diagnostic inaccuracies can diminish treatment effectiveness and increase the risk of adverse effects [1], [2]. Consequently, there is a need for clinical decision support systems (CDSS) capable of automatically identifying patterns between patient complaints and drug therapies without compromising transparency for medical professionals [3], [4].

The application of data mining in healthcare analysis has become increasingly important as medical data complexity grows, where the selection of appropriate methods significantly influences analytical performance [5]. Common approaches include classification and association rule mining algorithms. Within the classification category, Decision Tree (DT) is widely used due to its interpretability and ability to model relationships between attributes transparently, particularly in handling uncertainty in medical data [6], [7], [8], [9]. Its tree-based structure enables the

generation of decision rules that are easily understood by non-technical users, making it suitable for clinical decision support systems [10].

On the other hand, association rule mining (ARM) focuses on discovering correlation patterns between items that frequently co-occur within a dataset [11]. The ECLAT algorithm, which utilizes a vertical data representation, has proven efficient for frequent itemset mining in clinical data analysis [12]. ARM is particularly valuable in medical contexts because it can extract interpretable patterns directly from data and can complement supervised learning approaches [8], [13]. However, challenges arise when handling complex multi-label data, where capturing inter-label dependencies becomes more difficult [14]. In such cases, differences in algorithmic approaches can significantly affect performance outcomes, particularly in terms of precision and recall [15], [16].

Several previous studies have implemented classification algorithms for disease prediction [2], [8], while others have utilized association rule mining to explore correlation patterns between symptoms and medications [3], [13]. Although these approaches demonstrate significant potential, most existing studies focus on a single method and rarely perform comparative evaluations that simultaneously assess both rule characteristics and recommendation performance, particularly in multi-label clinical data. Comparative analysis between models is essential to ensure the stability and reliability of algorithmic performance across varying medical data characteristics [17]. Therefore, an in-depth study is required to analyze the behavioral differences between association-based (ECLAT) and classification-based (Decision Tree) algorithms in generating drug therapy recommendation rules, especially in terms of how their differing mechanisms influence both rule quality and recommendation effectiveness.

This condition indicates a persistent lack of studies directly comparing association rule mining and classification algorithms within the context of drug therapy recommendations based on complaints and diagnoses. These two approaches possess fundamentally different characteristics: ECLAT emphasizes discovering the strength of association patterns between items simultaneously, whereas Decision Tree focuses on predictive capability through a hierarchical decision-making process. As a result, both methods may produce significantly different rules even when applied to the same dataset. Without a systematic and quantitative comparison, particularly using recommendation-based evaluation metrics, the selection of methods for developing drug recommendation systems remains insufficiently evidence-based, making it difficult to determine the most appropriate model for clinical needs.

This issue becomes increasingly significant as drug recommendation systems are not only required to demonstrate high performance but must also be interpretable to gain acceptance from medical professionals. Systematic reviews indicate that although modern black-box machine learning models are growing in popularity, the lack of explainability remains a major barrier to clinical adoption; therefore, the use of transparent white-box models is recommended to ensure user trust [18]. This aligns with findings that rule-based approaches achieve higher acceptance levels because they allow clinicians to logically trace the basis of the system's decision-making process [13]. Consequently, an in-depth understanding of the characteristics of the rules generated by ECLAT and Decision Tree is crucial to ensure that the developed system is not only technically accurate but also relevant to practical healthcare needs.

Based on this background, this study aims to compare the performance of the ECLAT and Decision Tree algorithms in generating drug therapy recommendation rules based on patient complaints and diagnoses. This study hypothesizes that ECLAT will generate a larger number of rules with stronger association measures, while Decision Tree will produce fewer but more precise rules in terms of recommendation performance. The dataset used is sourced from patient visit records at a primary healthcare facility in Pangkep Regency. The comparison is conducted through experiments with variations in minimum support values and a comprehensive evaluation of the resulting rules. To ensure robust and unbiased evaluation, the experimental design employs cross-validation, where models are trained on a subset of the data and evaluated on unseen data. While rule quality is analyzed using various interest measures such as confidence, lift, leverage, conviction, and Zhang's metric [19], the primary focus of this study is the evaluation of recommendation performance using Top-K metrics, including Hit@k, Precision@k, and Recall@k. The integration of multiple interest measures remains important to filter misleading rules and improve rule validity [20], while correlation analysis is applied to measure inter-variable dependencies within the healthcare dataset [21].

Through this research, it is expected that a comprehensive understanding of the characteristics and performance of the ECLAT and Decision Tree algorithms in generating drug therapy recommendation rules will be obtained. This comparative approach is essential because algorithmic performance cannot be generalized across domains; rather, it is significantly influenced by data specifications and analytical objectives, as demonstrated in previous comparative studies of medical classification models [22]. Empirically, this study presents a structured comparison of both approaches using a unified experimental framework with cross-validation and evaluation based on recommendation performance metrics, including Hit@k, Precision@k, and Recall@k. The practical contribution of this research is intended to serve as an evidence-based foundation for clinical decision support system developers and medical professionals in selecting the most suitable method. Conceptually, the mapping of patterns from historical data to therapy recommendations reinforces the effectiveness of data-driven decision-making in the healthcare sector [23]. This study also contributes by providing a reproducible and fair evaluation framework for comparing rule-based and classification-based approaches in clinical recommendation systems.

2. Method

This study employs a comparative design to extract drug therapy recommendation rules through the integration of ECLAT and Decision Tree algorithms. The selection of a rule-based approach is based on its superior clinical stability compared to collaborative methods when applied to primary medical record data [1], [8]. Specifically, the use of Decision Tree (DT) provides an advantage in forming hierarchical decision structures that support clinical transparency [24]. The dataset used is sourced from the BPJS Kesehatan website, focusing on patient visit records at a public health center (Puskesmas) in Pangkajene and Kepulauan (Pangkep) Regency. Data were collected between August 25, 2025, and October 17, 2025, totaling 1,000 visit records. The primary attributes analyzed include patient complaints, diagnoses, and drug therapies.

To ensure a fair and unbiased evaluation, the experimental process employs a 5-fold cross-validation strategy. In each fold, the dataset is divided into training and testing subsets. Rule generation for both ECLAT and Decision Tree is performed exclusively on the training data, while recommendation performance is evaluated on unseen test data. This approach prevents data leakage and ensures that the evaluation reflects the generalization capability of the models [25]. A rule-based approach guarantees accountability through direct interpretation of the correlation between patient conditions and therapies [13], [26]. This aligns with the characteristics of DT, which is capable of generating systematic decision structures based on clinical attributes, ensuring that every recommendation can be explicitly traced and justified [27], [28]. The research workflow, which integrates the preprocessing stage, the parallel processing of both algorithms, and the comparative evaluation, is visualized in Figure 1.

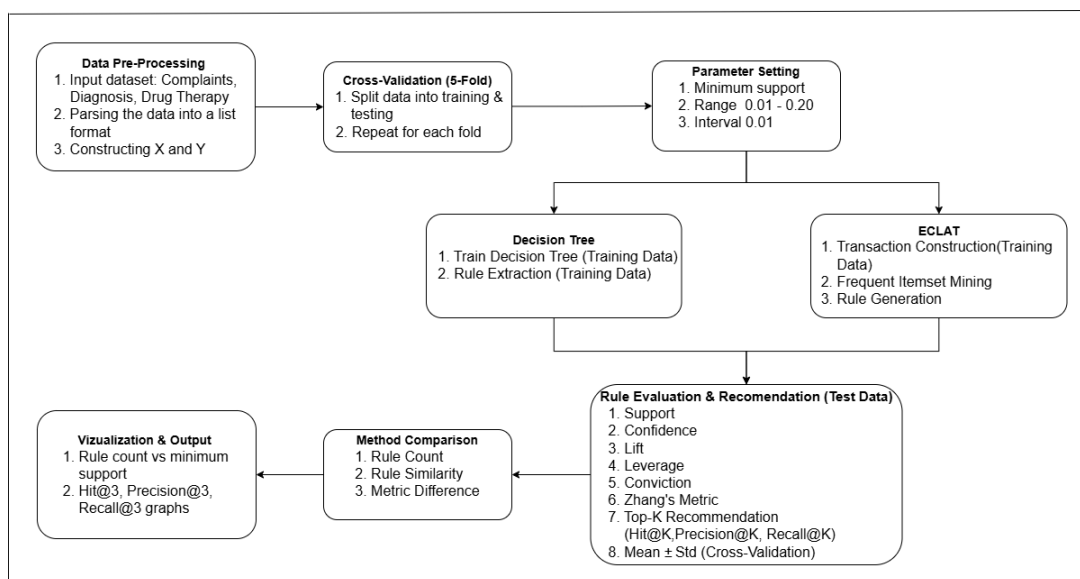


Figure 1. Proposed Research Workflow

Data Pre-Processing

The initial stage of the research begins with loading the dataset into the system and converting the values of the complaints, diagnoses, and drug therapy attributes into item lists. This process ensures that each value is treated as a discrete item that can be processed by association algorithms.

Subsequently, the complaints and diagnoses attributes are merged to form a set of patient condition features, which serve as antecedent candidates, while the drug therapy attribute serves as the consequent candidate. For Decision Tree processing, the combination of complaints and diagnoses is represented as a binary vector using the multi-hot encoding technique, whereas the drug therapy is represented as a binary target vector [29]. For ECLAT, each data row is structured into a transaction containing all complaint, diagnosis, and therapy items.

The output of this stage consists of two data representations: a transaction-based representation for ECLAT and a feature-based representation for Decision Tree, both of which are derived from the same data source.

Cross Validation (5-Fold)

To ensure a fair and unbiased evaluation, this study employs a 5-fold cross-validation strategy. The dataset is partitioned into five equal subsets, where in each iteration, four subsets are used as training data and one subset is used as testing data. This process is repeated five times so that each subset serves as the test data exactly once.

In each fold, rule generation for both ECLAT and Decision Tree is performed exclusively using the training data. The extracted rules are then applied to the corresponding test data to generate drug therapy recommendations. This approach ensures that the evaluation is conducted on unseen data, thereby preventing data leakage and providing a more reliable estimation of model generalization.

The performance metrics obtained from each fold are aggregated by computing the mean and standard deviation to reflect both the effectiveness and stability of the models across different data partitions.

Parameter Setting

The parameter setting is conducted to ensure that the rule generation process in both ECLAT and Decision Tree algorithms runs consistently and allows for a fair comparative analysis. All primary parameters are established before the algorithmic processing begins, as illustrated in the Parameter Setting stage within the research workflow in [Figure 1](#).

The minimum support parameter serves as the primary threshold for determining the frequency of patterns considered significant. The minimum support values are varied within a range of 0.01 to 0.20, with an interval of 0.01. This variation aims to observe the influence of changing frequency thresholds on the number of generated rules, rule characteristics, and the quality of drug recommendations.

The antecedent domain for both methods is restricted solely to the patient complaint and diagnosis attributes, while the consequent domain is limited to the drug therapy attribute. This domain restriction is implemented to ensure that the generated rules possess clear clinical interpretation, specifically mapping patient conditions to drug recommendations.

The maximum number of items in the consequent is set to three drug items for each rule. This restriction aims to manage rule complexity, ensuring they remain easy to interpret and align with common combination therapy patterns found within the data.

In the Decision Tree algorithm, the node splitting criterion utilizes the Gini index [30]. The tree depth is not restricted, allowing the model to flexibly extract patterns according to the data distribution. The parameter for the minimum number of samples at a leaf (minimum samples leaf) is determined by multiplying the minimum support value by the number of training samples in each fold. This approach ensures that the frequency threshold in the Decision Tree is aligned with the minimum support concept in ECLAT, enabling a fair comparison between both methods under the same frequency constraint. A summary of the parameters used for both algorithms is presented in [Table 1](#).

Table 1. Parameter Setting

Parameter	ECLAT	Decision Tree
Minimum Support	0.01-0.20 (step 0.01)	0.01-0.20 (step 0.01)
Domain Antecedent	Complaints, Diagnosis	Complaints, Diagnosis
Domain Consequent	Drug therapy	Drug therapy
Max Consequent Items	3	3
Splitting Criterion	-	Gini
Max Depth	-	Unrestricted
Min Sample Leaf	-	Support x training data (per fold)

Eclat Processing

For each minimum support value and within each fold of the cross-validation process, the system constructs a vertical data representation using only the training data, where each item is associated with a list of transactions in which it appears. Based on this representation, frequent itemsets are identified by combining items that satisfy the minimum support threshold.

The frequent itemsets are subsequently filtered to form association rules with an antecedent $\rightarrow$ consequent structure, where the antecedent is derived from complaints and diagnoses, while the consequent is derived from drug therapies. For each generated rule, association metrics are calculated to assess the strength of the relationships between items.

The extracted rules are then applied to the corresponding test data to generate drug therapy recommendations, ensuring that the evaluation is performed on unseen data. All rules, along with their respective metric values, are stored in Excel spreadsheets according to the minimum support values used.

Decision Tree Processing

Within each fold of the cross-validation process, the dataset that has undergone the multi-hot encoding process is used to train the Decision Tree model using only the training data. Each minimum support value is converted into a minimum samples leaf parameter based on the size of the training data, ensuring that the generated tree only considers patterns that appear at least as frequently as the specified support threshold.

Following the training process, the tree structure is extracted into a set of IF–THEN rules. This approach is commonly employed in medical decision support systems because the rules generated from Decision Trees are explicit, easily traceable, and can be validated by medical professionals [2]. These rules are then normalized into an antecedent $\rightarrow$ consequent format to ensure comparability with the rules generated by ECLAT. Each rule's association metric values are subsequently calculated using the same formulas.

The extracted rules are applied to the corresponding test data in each fold to generate drug therapy recommendations, ensuring that evaluation is performed on unseen data. The rules and metric summaries are stored in Excel format.

Rule Evaluation Metric Formulation

The evaluation of association rule quality is conducted using several statistical metrics commonly employed in association rule mining, namely support, confidence, lift, leverage, conviction, and Zhang's metric. Support represents the co-occurrence frequency of the antecedent and consequent within the training data in each fold. Confidence indicates the probability of the consequent appearing when the antecedent is present. Lift is used to measure the degree of dependence between the antecedent and consequent, where a lift value greater than one indicates a positive relationship. Leverage shows the difference between the observed probability of co-occurrence and the expected probability if the antecedent and consequent were independent. Conviction describes the implication strength of a rule

by considering the prediction failure rate. Meanwhile, Zhang's metric is utilized to assess the strength of the association while accounting for imbalances in data distribution.

Suppose there are two sets of items, namely antecedent X and consequent Y , with the total number of transactions in the training data denoted as N . The probability of occurrence of an itemset is expressed as the proportion of transactions containing that itemset relative to all transactions, defined as $P(X)$, $P(Y)$, and $P(XY)$. All probability values range between 0 and 1 and serve as the foundation for selecting the best association rules. All rule evaluation metrics are formulated as follows:

$$Support(X, Y) = \frac{|T(X \cap Y)|}{N} \quad (1)$$

$$Confidence(X \rightarrow Y) = \frac{P(XY)}{P(X)} \quad (2)$$

$$Lift(X \rightarrow Y) = \frac{Confidence(X \rightarrow Y)}{P(Y)} \quad (3)$$

$$Leverage(X \rightarrow Y) = P(X, Y) - P(X) \times P(Y) \quad (4)$$

$$Conviction(X \rightarrow Y) = \frac{1 - P(Y)}{1 - P(X \rightarrow Y)} \quad (5)$$

$$A = P(XY) - P(X)P(Y) \quad (6.a)$$

$$B = \max\{P(XY)(1 - P(X)), P(X)(P(Y) - P(XY))\} \quad (6.b)$$

$$Zhang(X \rightarrow Y) = \frac{A}{B} \quad (6.c)$$

Rule Evaluation and Top-K Recommendation

The association rules generated by each method serve as the basis for the drug recommendation process. Within each fold of the cross-validation process, the rules generated from the training data are applied to the corresponding test data. For each patient record, the patient's condition is matched against the rules' antecedents. If a match is found, the consequent of that rule is designated as a recommendation candidate. All recommendation candidates are then ranked based on the highest confidence values, and the top K recommendations are selected as the Top-K set.

Since each patient record may contain multiple drug therapies, the evaluation is conducted in a multi-label setting. A recommendation is considered correct if at least one of the predicted drugs matches any of the actual drugs in the test data.

The performance evaluation of the recommendation system is conducted using Hit@K, Precision@K, and Recall@K metrics. Hit@K is used to measure the system's success in generating at least one correct recommendation within the Top-K. Precision@K measures the proportion of relevant recommendations out of all provided recommendations. Recall@K measures the system's ability to identify the correct drugs out of all actual drugs prescribed. The Top-K evaluation metrics are formulated as follows:

$$Hit@K = \frac{1}{M} \sum_{i=1}^M h_i \quad (7)$$

$$Precision@K = \frac{|R_k \cap T|}{K} \quad (8)$$

$$Recall@K = \frac{|R_k \cap T|}{|T|} \quad (9)$$

Where $h_i = 1$ if at least one correct drug appears in the Top-K recommendations for the i -th data record, and $h_i = 0$ otherwise. Furthermore, R_k represents the set of Top-K recommended drugs, T denotes the set of actual drugs, and M signifies the total number of test data records. All evaluation results are reported as the mean and standard deviation across the cross-validation folds to reflect both performance and stability.

Method Comparison

The summarized results of ECLAT and Decision Tree are aggregated across cross-validation folds and merged based on their minimum support values. The comparison is conducted in terms of the number of generated rules, the average values of association metrics, and the recommendation performance metrics. All performance results are reported as mean and standard deviation to reflect both effectiveness and stability across different data partitions.

In addition, a rule similarity analysis is performed by calculating the number of identical rules, rules generated exclusively by ECLAT, and rules generated exclusively by the Decision Tree. This analysis is conducted using a representative fold to examine structural differences between rule sets. Differences in confidence values, Hit@3, Precision@3, and Recall@3 are also calculated as indicators of performance variation between the two methods.

Visualization and Output

The research findings are visualized through a series of comparative charts to provide a comprehensive analysis of the model's behavior. These include graphs depicting the relationship between the number of generated rules and the minimum support threshold, as well as visualizations of recommendation performance using Hit@3, Precision@3, and Recall@3 metrics. The performance graphs are presented with error bars representing the mean and standard deviation obtained from cross-validation, enabling a clearer interpretation of both effectiveness and stability across different parameter settings.

To ensure data integrity and facilitate further analysis, all numerical results, performance metrics, and graphical representations are consolidated into a single structured Excel file. This unified output serves as the final report, encompassing the complete experimental results. This structured documentation is designed to support clinical validation by medical professionals and provide a reproducible baseline for future drug recommendation research.

3. Result & Discussion

Results

This section presents the experimental results and comparative analysis between the ECLAT and Decision Tree algorithms in extracting association rules and generating drug recommendations based on patient visit data. The results discussed encompass the number of generated rules, rule quality based on association metrics, and the performance of the recommendation system using the Top-K approach. All results are obtained using a cross-validation strategy, where rule generation is performed on training data and evaluation is conducted on unseen test data. The reported performance values are presented as mean and standard deviation to reflect both effectiveness and stability. Each key finding is interpreted to address the research objectives, namely assessing the completeness, quality, and effectiveness of the recommendations from both methods.

Comparison of Rule Quantity and Completeness

Experimental results consistently show that the ECLAT algorithm generates a higher number of rules compared to the Decision Tree across all tested minimum support values. This pattern indicates that ECLAT possesses a more comprehensive search space exploration capability, enabling it to capture a wider range of relationship combinations between complaints, diagnoses, and drug therapies. By uncovering these intricate patterns, ECLAT provides a more granular view of medical associations that might otherwise remain hidden in more restrictive models. This finding is consistent with previous research stating that association rule mining tends to produce richer knowledge, yet requires further selection processes for clinical implementation needs [8], [13].

Table 2. Comparison of Number of Rules Generated by ECLAT and Decision Tree

Min Support	Total Rules		Intersecting Rules	Exclusive Rules	
	ECLAT	DT		ECLAT	DT
0.01	1298	795	795	503	0
0.02	498	279	279	219	0
0.03	282	159	159	123	0
0.04	166	97	97	69	0
0.05	113	65	65	48	0
0.06	73	37	37	36	0
0.07	54	23	23	31	0
0.08	41	17	17	24	0
0.09	28	16	16	12	0
0.10	25	13	13	12	0
0.11	18	10	10	8	0
0.12	14	7	7	7	0
0.13	10	5	5	5	0
0.14	7	5	5	2	0
0.15	5	3	3	2	0
0.16	3	2	2	1	0
0.17	3	2	2	1	0
0.18	3	2	2	1	0
0.19	3	2	2	1	0
0.20	1	1	1	0	0

The reported number of rules represents the average results across cross-validation folds, ensuring that the observed patterns are consistent across different data partitions.

Based on [Table 2](#), the DT Exclusive Rule column shows a value of zero across all minimum support levels, indicating that no rules are generated exclusively by the Decision Tree. All rules obtained by the Decision Tree are also found within the rule set produced by ECLAT under the same experimental conditions. This observation is validated through rule set comparison and suggests that the Decision Tree effectively selects a subset of patterns identified by ECLAT. This behavior can be attributed to the Decision Tree mechanism, which prioritizes splits that maximize information gain or minimize impurity, thereby filtering out less relevant associations. In contrast, ECLAT's exhaustive search approach allows it to explore a broader range of item combinations without relying on heuristic pruning.

At low minimum support values, the discrepancy in the number of rules between the two algorithms is significant. For instance, at a minimum support of 0.01, ECLAT generates over a thousand rules, whereas the Decision Tree yields only a few hundred. This gap reflects that many low-frequency yet valid patterns are successfully captured by ECLAT but are excluded by the Decision Tree due to its tree construction mechanism and pruning process. While pruning is beneficial for preventing overfitting, it inherently reduces the granularity of niche medical patterns. Consequently, ECLAT is more suitable for exploratory data mining scenarios where rare disease–drug correlations are of interest.

As the minimum support value increases, the number of rules generated by both methods decreases drastically. In the medium range of minimum support, the discrepancy in the number of rules begins to narrow, and at high minimum support levels, both algorithms yield an identical and very small number of rules. This condition indicates a convergence of results when only highly frequent patterns are retained.

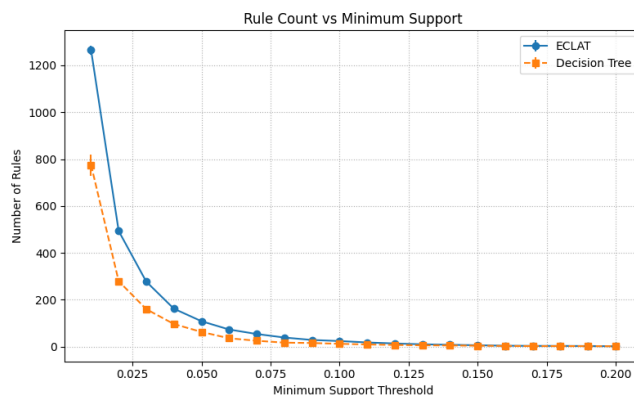


Figure 2. Number of Rules versus Minimum Support

The graph in [Figure 2](#) illustrates a sharp downward trend in the number of rules as minimum support increases, as well as the dominance of ECLAT at low to medium support levels. This finding addresses the research question regarding the completeness of the methods, where ECLAT demonstrates a stronger capability in discovering a broader set of patterns that meet the support criteria compared to the Decision Tree. This characteristic is aligned with the study on the application of ECLAT for diabetes prediction by Kumari [10], which found that this algorithm is highly effective in generating complete frequent itemsets without candidate generation, thereby enabling the identification of potentially important clinical patterns in complex medical datasets.

Comparison of Rule Quality and Recommendation Performance between ECLAT and Decision Tree

In addition to analyzing the quantity and completeness of the rules generated by each method, this study evaluates the quality of the association rules and their impact on the performance of the drug recommendation system. Rule quality evaluation is conducted using association metrics, including confidence and lift as primary indicators of the relationship strength between items. Meanwhile, recommendation performance is evaluated using Top-K recommendation metrics, specifically Hit@3, Precision@3, and Recall@3, which represent the system's success in suggesting drugs relevant to the patient's condition. To provide a general overview of both methods' performance across various minimum support scenarios, the results are aggregated across cross-validation folds and reported as mean and standard deviation values.

Table 3. Average Performance Comparison between ECLAT and Decision Tree

Metric	ECLAT	Decision Tree
Average Confidence ↑	0.5783	0.4997
Average Lift ↑	2.2070	2.1518
Precision@3 ↑	0.3048 ± 0.0200	0.2652 ± 0.0196
Recall@3 ↑	0.2417 ± 0.0136	0.2069 ± 0.0153
Hit@3 ↑	0.5510 ± 0.0325	0.4674 ± 0.0291

Table 3 presents a comparison of the average performance between ECLAT and the Decision Tree based on rule quality and recommendation metrics. All reported values are averaged across cross-validation folds, with standard deviation indicating performance stability. The results indicate that ECLAT achieves higher average confidence (0.5783) and average lift (2.2070) values compared to the Decision Tree (0.4997 and 2.1518, respectively). The lift values greater than one confirm the presence of positive dependencies between drug items and diagnoses, indicating that certain therapies are more likely to be prescribed under specific clinical conditions.

This finding supports the argument by Akbaş et al. [19], which states that interest measures such as lift and conviction are essential for validating the strength of medical rules, as high confidence alone may be misleading without sufficient lift. In this context, the higher lift values observed in ECLAT suggest stronger and more reliable

associations. This behavior aligns with ECLAT's characteristics as an exhaustive pattern mining algorithm, enabling it to explore a broader combination of itemsets and capture more detailed relationships between complaints, diagnoses, and drug therapies. Consequently, ECLAT can be considered a more comprehensive approach for knowledge discovery, particularly in identifying less frequent but clinically meaningful associations.

In terms of recommendation performance, ECLAT also achieves higher Precision@3 (0.3048 ± 0.0200) and Recall@3 (0.2417 ± 0.0136) compared to the Decision Tree (0.2652 ± 0.0196 and 0.2069 ± 0.0153). This indicates that ECLAT produces more relevant recommendations overall. However, the Decision Tree demonstrates more stable performance, as indicated by slightly lower standard deviation values across folds. This observation is consistent with the findings of Sae-Ang et al. [1], which suggest that classification-based approaches tend to produce more consistent predictions due to their direct mapping between input features and target labels.

Furthermore, regarding the Hit@3 metric, ECLAT achieves higher values (0.5510 ± 0.0325) compared to the Decision Tree (0.4674 ± 0.0291), indicating a higher probability of generating at least one correct recommendation within the Top-K results. This reflects ECLAT's strength in coverage, where a larger set of rules increases the likelihood of matching relevant patterns in patient data.

These contrasting behaviors are further illustrated in **Figure 3 – 5**, where ECLAT consistently maintains higher recall and hit values across various minimum support thresholds, while the Decision Tree shows more stable trends with smaller variability. Overall, these results highlight a trade-off between pattern completeness and recommendation stability: ECLAT excels in discovering comprehensive and strong association patterns, whereas the Decision Tree provides more stable and structured recommendations with fewer rules.

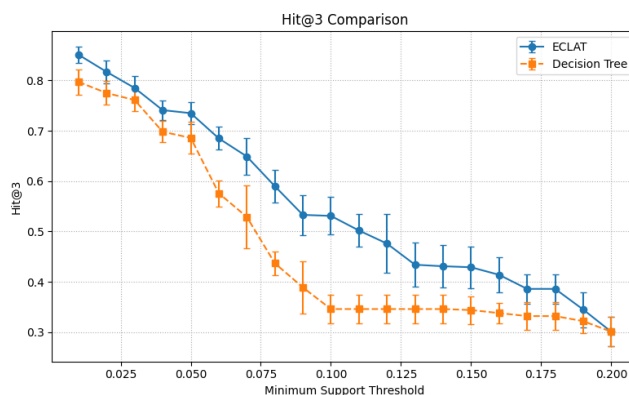


Figure 3. Hit@3 versus Minimum Support

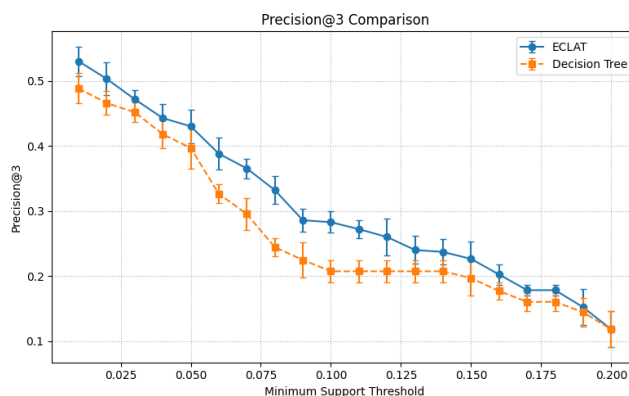


Figure 4. Precision@3 versus Minimum Support

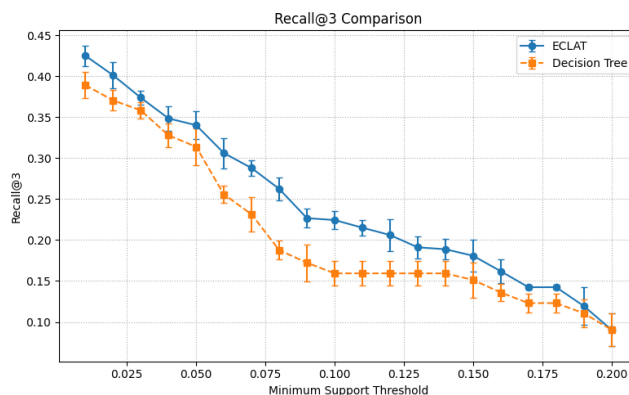


Figure 5. Recall@3 versus Minimum Support

Figure 3 – 5 illustrate the comparison of Hit@3, Precision@3, and Recall@3 across different minimum support values. The reported values represent the average performance across cross-validation folds, with error bars indicating standard deviation.

Overall, ECLAT consistently achieves higher values for Hit@3, Precision@3, and Recall@3 compared to the Decision Tree across most minimum support thresholds. This indicates that ECLAT is able to generate more relevant and comprehensive recommendations by leveraging a larger set of association rules. The difference between the two methods is more pronounced at lower minimum support values, where ECLAT benefits from a broader exploration of pattern combinations.

As the minimum support value increases, the performance of both methods gradually decreases and tends to converge. This trend is consistent with the reduction in the number of generated rules, as higher support thresholds restrict the pattern search space to only the most frequent itemsets. Consequently, both methods rely on a smaller and more similar set of patterns at higher thresholds.

Although ECLAT demonstrates higher overall performance, the Decision Tree exhibits more stable behavior, as indicated by its relatively smaller variability across support values. This reflects the inherent characteristics of classification-based models, which prioritize structured decision boundaries over exhaustive pattern exploration.

Based on these results, it can be concluded that ECLAT is more effective in capturing comprehensive association patterns and achieving higher recommendation coverage, while the Decision Tree provides more stable and structured recommendations. This highlights a trade-off between pattern completeness and model stability in the development of drug recommendation systems.

Discussion

The findings of this study highlight a clear contrast between the exploratory strength of association rule mining and the structured selectivity of classification-based approaches in the context of drug therapy recommendation systems. The experimental results, which are averaged across cross-validation folds, confirm that ECLAT and Decision Tree exhibit fundamentally different behaviors that directly influence the nature and usability of the generated rules.

From a knowledge discovery perspective, ECLAT demonstrates a strong ability to uncover comprehensive association patterns. This is evident from the significantly higher number of generated rules across all minimum support values. Such behavior indicates that ECLAT is highly effective in identifying both dominant and less frequent relationships between patient complaints, diagnoses, and drug therapies. This characteristic is particularly important in the medical domain, where less frequent patterns may still hold clinical significance, especially in cases involving uncommon disease combinations or personalized treatments. However, this exhaustive pattern generation may also introduce redundancy and potential noise, requiring additional filtering mechanisms before practical implementation.

In contrast, the Decision Tree algorithm produces a more compact and structured set of rules. The results show that all rules generated by the Decision Tree are subsets of those produced by ECLAT, indicating that Decision Tree inherently performs a selection process during model construction. By prioritizing splits based on impurity reduction, Decision Tree effectively filters out weaker or less relevant associations. This built-in pruning mechanism leads to fewer but more focused rules, which contributes to more stable recommendation behavior across cross-validation folds.

The evaluation of rule quality further reinforces these differences. ECLAT achieves higher average confidence and lift values, suggesting that the associations it discovers are statistically stronger and more informative. This implies that ECLAT is more capable of capturing meaningful relationships beyond random co-occurrence. In addition, ECLAT also achieves higher Precision@3, Recall@3, and Hit@3 values compared to the Decision Tree, indicating better overall recommendation performance in terms of relevance and coverage.

However, despite lower average performance values, the Decision Tree demonstrates relatively smaller standard deviations across folds, indicating more consistent and stable predictions. This reflects the inherent characteristics of supervised learning models, which optimize decision boundaries based on labeled outcomes and tend to generalize more consistently across different data partitions.

Another important observation is the convergence behavior at higher minimum support values. As the support threshold increases, both algorithms generate fewer rules and exhibit increasingly similar performance. This suggests that when only highly frequent patterns are considered, the distinction between exhaustive and selective approaches becomes less significant. In practical terms, this implies that the choice of algorithm becomes more critical when dealing with low-support scenarios, where hidden or rare patterns are of interest.

Overall, the results reveal a trade-off between completeness and stability. ECLAT excels in generating a comprehensive set of association rules and achieving higher recommendation coverage, making it highly suitable for exploratory analysis and knowledge discovery in clinical datasets. On the other hand, Decision Tree provides more stable and structured recommendations with fewer rules, making it more appropriate for real-world decision support systems where consistency and interpretability are prioritized.

These findings suggest that the two approaches are not mutually exclusive but rather complementary. A potential future direction is the integration of both methods, where ECLAT is used to generate a rich pool of candidate rules, and Decision Tree or other classification techniques are employed to refine and prioritize these rules for recommendation purposes. Such a hybrid approach could leverage the strengths of both algorithms, resulting in a more robust and effective clinical decision support system.

4. Conclusion

This study compares the ECLAT and Decision Tree algorithms in generating drug therapy recommendation rules based on patient complaints, diagnoses, and therapy data. The experimental results, obtained through cross-validation, indicate that ECLAT produces a more comprehensive set of association rules with higher average confidence and lift values, reflecting stronger pattern discovery capability. In addition, ECLAT achieves higher Precision@3, Recall@3, and Hit@3 values, indicating better overall recommendation performance in terms of relevance and coverage.

In contrast, the Decision Tree generates fewer but more structured rules and demonstrates more stable performance across folds, as indicated by lower variability in evaluation metrics. These findings highlight a trade-off between pattern completeness and model stability, where ECLAT is more suitable for knowledge exploration and comprehensive pattern discovery, whereas the Decision Tree is more appropriate for applications requiring consistent and interpretable recommendations.

This study contributes a systematic comparative analysis using real-world clinical data and provides insights for developing transparent and interpretable clinical decision support systems. Future work may focus on hybrid approaches that integrate both methods and incorporate more diverse datasets to further improve recommendation performance and robustness.

ACKNOWLEDGMENT

The authors would like to express their sincere gratitude to the Community Health Center in Pangkajene and Kepulauan Regency (Pangkep) for the permission, support, and cooperation provided during the process of collecting and providing research data. This support greatly facilitated the implementation of the study, particularly in obtaining relevant and structured medical record data, enabling this research to be completed successfully.

References:

- [1] A. Sae-Ang, S. Chairat, N. Tansuebchueasai, O. Fumaneechoat, T. Ingviya, and S. Chaichulee, "Drug Recommendation from Diagnosis Codes: Classification vs. Collaborative Filtering Approaches," *Int. J. Environ. Res. Public Health*, vol. 20, no. 1, Jan. 2023, doi: <https://doi.org/10.3390/ijerph20010309>.
- [2] A. S. Ahmed and H. A. Salah, "A comparative study of classification techniques in data mining algorithms used for medical diagnosis based on DSS," *Bulletin of Electrical Engineering and Informatics*, vol. 12, no. 5, pp. 2964–2977, Oct. 2023, doi: <https://doi.org/10.11591/eei.v12i5.4804>.
- [3] A. Mai, K. Voigt, J. Schübel, and F. Gräßer, "A drug recommender system for the treatment of hypertension," *BMC Med. Inform. Decis. Mak.*, vol. 23, no. 1, Dec. 2023, doi: <https://doi.org/10.1186/s12911-023-02170-y>.
- [4] P. Mateos and A. Bellogín, "A systematic literature review of recent advances on context-aware recommender systems," *Artif. Intell. Rev.*, vol. 58, no. 1, Jan. 2025, doi: <https://doi.org/10.1007/s10462-024-10939-4>.
- [5] A. Muh. F. Asfar, M. Hasnawi, and H. Darwis, "Combinations of Feature Extractions and Machine Learning Algorithms for Skin Cancer Classification," *Jurnal Teknik Informatika (Jutif)*, vol. 5, no. 6, pp. 1591–1598, Dec. 2024, doi: <https://doi.org/10.52436/1.jutif.2024.5.6.2514>.
- [6] N. K. and M. K. M. B., "Fuzzy rule based classifier model for evidence based clinical decision support systems," *Intelligent Systems with Applications*, vol. 22, Jun. 2024, doi: <https://doi.org/10.1016/j.iswa.2024.200393>.
- [7] V. S. K. Reddy, P. Meghana, N. V. S. Reddy, and B. A. Rao, "Prediction on Cardiovascular disease using Decision tree and Naïve Bayes classifiers," in *Journal of Physics: Conference Series*, IOP Publishing Ltd, Jan. 2022. doi: <https://doi.org/10.1088/1742-6596/2161/1/012015>.
- [8] D. Wendimu and K. Biredagn, "Developing a knowledge-based system for diagnosis and treatment recommendation of neonatal diseases," *Cogent Eng.*, vol. 10, no. 1, 2023, doi: <https://doi.org/10.1080/23311916.2022.2153567>.
- [9] Muh. I. E. Saputra Troy, S. R. Jabir, and S. Anraeni, "Evaluation of Multi-Class Classification Performance Lung Cancer Through K-NN and SVM Approach," *ILKOM Jurnal Ilmiah*, vol. 17, no. 1, pp. 27–33, Apr. 2025, doi: <https://doi.org/10.33096/ilkom.v17i1.2464.27-33>.
- [10] S. Zahoor, P. Liò, G. Dias, and M. Hasanuzzaman, "Integrating probabilistic trees and causal networks for clinical and epidemiological data," *Artif. Intell. Med.*, vol. 173, Mar. 2026, doi: <https://doi.org/10.1016/j.artmed.2026.103350>.
- [11] J. Cui, S. Zhao, and X. Sun, "An Association Rule Mining Algorithm for Clinical Decision Support," in *ACM International Conference Proceeding Series*, Association for Computing Machinery, Mar. 2022, pp. 137–143. doi: <https://doi.org/10.1145/3532213.3532234>.
- [12] P. S. Kumari, "Predicting the Severity of Diabetes Using ECLAT Algorithm in Data Mining," 2024, pp. 359–370. doi: https://doi.org/10.2991/978-94-6463-433-4_26.
- [13] G. T. Berge, O. C. Granmo, T. O. Tveit, A. L. Ruthjersen, and J. Sharma, "Combining unsupervised, supervised and rule-based learning: the case of detecting patient allergies in electronic health records," *BMC Med. Inform. Decis. Mak.*, vol. 23, no. 1, Dec. 2023, doi: <https://doi.org/10.1186/s12911-023-02271-8>.

- [14] R. Alazaidah, G. Samara, S. Almatarneh, M. Hassan, M. Aljaidi, and H. Mansur, "Multi-Label Classification Based on Associations," *Applied Sciences (Switzerland)*, vol. 13, no. 8, Apr. 2023, doi: <https://doi.org/10.3390/app13085081>.
- [15] P. Lestari, L. Belluano, R. A. Rahma, H. Darwis, and A. R. Manga, "Analysis of ensemble machine learning classification comparison on the skin cancer MNIST dataset," *Computer Science and Information Technologies*, vol. 5, no. 3, pp. 235–242, 2024, doi: <https://doi.org/10.11591/csit.v5i3.pp235-242>.
- [16] S. Fi. Nurul Fitri H, F. Fattah, and H. Azis, "Comparative Analysis of Machine Learning Algorithm Variations in Classifying Body Shaming Topics on Social Media X," *Indonesian Journal of Data and Science*, vol. 5, no. 2, Jul. 2024, doi: <https://doi.org/10.56705/ijodas.v5i2.82>.
- [17] A. Rachman Manga, A. Putri Utami, H. Azis, Y. Salim, and A. Faradibah, "Optimizing classification models for medical image diagnosis: a comparative analysis on multi-class datasets," *Computer Science and Information Technologies*, vol. 5, no. 3, pp. 205–214, 2024, doi: <https://doi.org/10.11591/csit.v5i3.pp205-214>.
- [18] Y. Du, C. McNestry, L. Wei, A. M. Antoniadis, F. M. McAuliffe, and C. Mooney, "Machine learning-based clinical decision support systems for pregnancy care: A systematic review," May 01, 2023, *Elsevier Ireland Ltd.* doi: <https://doi.org/10.1016/j.ijmedinf.2023.105040>.
- [19] K. E. AKBAŞ *et al.*, "Assessment of Association Rule Mining Using Interest Measures on the Gene Data," *Medical Records*, vol. 4, no. 3, pp. 286–292, Sep. 2022, doi: <https://doi.org/10.37990/medr.1088631>.
- [20] K. Kelesidis, N. Fotopoulou, and D. Dervos, "Correlation as an ARM Interestingness Measure for Numeric Datasets," in *ACM International Conference Proceeding Series*, Association for Computing Machinery, Nov. 2023, pp. 1–7. doi: <https://doi.org/10.1145/3635059.3635060>.
- [21] L. F. G. Morales, P. Valdiviezo-Diaz, R. Reátegui, and L. Barba-Guaman, "Drug Recommendation System for Diabetes Using a Collaborative Filtering and Clustering Approach: Development and Performance Evaluation," *J. Med. Internet Res.*, vol. 24, no. 7, Jul. 2022, doi: <https://doi.org/10.2196/37233>.
- [22] Amaliah Faradibah, Dewi Widyawati, A. Ulfah Tenripada Syahar, and Sitti Rahmah Jabir, "Comparison Analysis of Random Forest Classifier, Support Vector Machine, and Artificial Neural Network Performance in Multiclass Brain Tumor Classification," *Indonesian Journal of Data and Science*, vol. 4, no. 2, pp. 54–63, Jul. 2023, doi: <https://doi.org/10.56705/ijodas.v4i2.73>.
- [23] H. Darwis, F. A. Syahrir, and L. N. Hayati, "A Hybrid Movie Recommendation System to Address Data Sparsity Using Genre-Based K-Means and Neural Collaborative Filtering," *ILKOM Jurnal Ilmiah*, vol. 17, no. 2, pp. 203–212, Sep. 2025, doi: <https://doi.org/10.33096/ilkom.v17i2.2868.203-212>.
- [24] Sitti Rahmah Jabir, Huzain Azis, and St. Hajrah Mansyur, "Enhancing The Quality of College Decisions Through Decision Tree and Random Forest Models," *Journal of Embedded Systems, Security and Intelligent Systems*, vol. 5, no. 1, pp. 13–18, Mar. 2024, doi: <https://doi.org/10.59562/jessi.v5i1.1225>.
- [25] Y. Peng *et al.*, "Uncertainty-Aware Explainable Recommendation with Large Language Models," Jan. 2024, <http://arxiv.org/abs/2402.03366>.
- [26] L. Verboven *et al.*, "A treatment recommender clinical decision support system for personalized medicine: method development and proof-of-concept for drug resistant tuberculosis," *BMC Med. Inform. Decis. Mak.*, vol. 22, no. 1, Dec. 2022, doi: <https://doi.org/10.1186/s12911-022-01790-0>.
- [27] Dewi Widyawati and Amaliah Faradibah, "Comparison Analysis of Classification Model Performance in Lung Cancer Prediction Using Decision Tree, Naive Bayes, and Support Vector Machine," *Indonesian Journal of Data and Science*, vol. 4, no. 2, pp. 80–89, Jul. 2023, doi: <https://doi.org/10.56705/ijodas.v4i2.76>.

- [28] H. Azis and S. R. Jabir, “Chemical Composition and Aroma Profiling: Decision Tree Modeling of Formalin Tofu,” *Journal of Embedded Systems, Security and Intelligent Systems*, vol. 4, no. 2, pp. 206–211, Nov. 2023, doi: <https://doi.org/10.59562/jessi.v4i2.1162>.
- [29] H. Bae, B. Kang, and C. E. Kim, “Understanding clinical decision-making in traditional East Asian medicine through dimensionality reduction: An empirical investigation,” *Comput. Biol. Med.*, vol. 197, Oct. 2025, doi: <https://doi.org/10.1016/j.compbiomed.2025.111081>.
- [30] J. Chen, L. Wu, K. Liu, Y. Xu, S. He, and X. Bo, “EDST: a decision stump based ensemble algorithm for synergistic drug combination prediction,” *BMC Bioinformatics*, vol. 24, no. 1, Dec. 2023, doi: <https://doi.org/10.1186/s12859-023-05453-3>.