

Exploring multi-objective approaches for rule generation in rule-based machine learning

David von Proeck, Jonathan Wurth, Jörg Hähner, Michael Heider

Angaben zur Veröffentlichung / Publication details:

Proeck, David von, Jonathan Wurth, Jörg Hähner, and Michael Heider. 2026. "Exploring multi-objective approaches for rule generation in rule-based machine learning." In *GECCO '26 companion: proceedings of the Genetic and Evolutionary Computation Conference Companion, July 13-17, 2026, San Jose, Costa Rica*, edited by Leonardo Trujillo and Ting Hu, 1220–27. New York, NY: ACM. <https://doi.org/10.1145/3795101.3814669>.

Nutzungsbedingungen / Terms of use:

CC BY 4.0



Exploring Multi-objective Approaches for Rule Generation in Rule-based Machine Learning

David von Proeck
david.von.proeck-zvlcil@uni-a.de
Universität Augsburg
Augsburg, Germany

Jörg Hähner
joerg.haehner@uni-a.de
Universität Augsburg
Augsburg, Germany

Jonathan Wurth
jonathan.wurth@uni-a.de
Universität Augsburg
Augsburg, Germany

Michael Heider
michael.heider@uni-a.de
Universität Augsburg
Augsburg, Germany

Abstract

Critical systems increasingly demand continuously learning, self-adaptive AI models that remain transparent to human stakeholders. Rule-based machine learning offers a promising path due to its human-interpretable structure and flexibility. However, a key challenge persists: more general rules inherently reduce rule set complexity and thus enhance interpretability, yet high predictive performance is essential and overly general rules can reduce that. Currently, most existing rule discovery methods optimize only (relative) performance or merge generality into a single scalarized objective. To address this, we propose a Pareto-front approximating rule discovery approach leveraging NSGA-II. We investigate five distinct objectives to optimize against predictive performance. Integrated into the SupRB algorithm, our method is evaluated on four real-world regression tasks. Results demonstrate that especially the optimization of the novelty of rules (relative to the samples matched by each rule) can successfully improve overall rule set errors and complexity. These improvements are extensively validated through Bayesian statistical analysis, quantifying both effect sizes and posterior probabilities. Our work provides a key next step towards balancing performance and interpretability in critical AI systems.

CCS Concepts

• **Applied computing** → Multi-criterion optimization and decision-making; • **Computing methodologies** → **Rule learning**; *Learning linear models*; **Supervised learning by regression**; • **Theory of computation** → *Evolutionary algorithms*.

Keywords

Explainable AI, Rule-based Machine Learning, Learning Classifier Systems, Novelty Search, Evolutionary Computation

ACM Reference Format:

David von Proeck, Jonathan Wurth, Jörg Hähner, and Michael Heider. 2026. Exploring Multi-objective Approaches for Rule Generation in Rule-based

Machine Learning. In *Genetic and Evolutionary Computation Conference (GECCO Companion '26)*, July 13–17, 2026, San Jose, Costa Rica. ACM, New York, NY, USA, 8 pages. <https://doi.org/10.1145/3795101.3814669>

1 Introduction

Rule-based machine learning models derive their predictive power from the quality of the rules extracted during training. Conventional methods, especially evolutionary rule-based approaches [31, 30, 15, 25], optimize these rules using a scalar objective function. Often this can even include a weighted sum of objectives such as accuracy, generality, and complexity rather than purely focussing on accuracy [14]. While computationally efficient, scalarized approaches enforce a fixed trade-off across the input space, potentially producing structurally homogeneous rules that may overfit local patterns or fail to capture nuanced relationships in complex regression tasks. For any optimization approach, identifying the minimal rule set required for an optimal model [5] is difficult [22, 23]. A more practical strategy involves generating a diverse, high-quality rule pool, of which an optimal subset can be selected in a second step, increasing the likelihood that the final model will include the necessary combinations. The SupRB system explicitly separates these steps in its Rule Discovery (RD) and Solution Composition (SC) components, which in turn allows for the straightforward exchange of the RD component.

This paper introduces a Multi-Objective Rule Discovery (MO-RD) approach, replacing the scalar fitness function with Pareto-optimality via NSGA-II [11]. By optimizing objectives such as Mean Squared Error (MSE) and novelty simultaneously, the method maintains dynamic trade-offs and preserves diversity through the use of NSGA-II's crowding distance mechanism. We evaluate the impact of secondary objectives, specifically generality and diversity metrics, on four real-world regression problems.

Replacing scalar fitness functions with a Pareto-based discovery process in rule discovery yields several key advantages, more specifically:

- Preserved predictive performance: transitioning to a Pareto-based process does not lead to a degradation in predictive performance.
- Greater model interpretability: MO-RD facilitates the construction of more compact models, increasing explainability.



This work is licensed under a Creative Commons Attribution 4.0 International License.
GECCO Companion '26, San Jose, Costa Rica
© 2026 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-2488-6/2026/07
<https://doi.org/10.1145/3795101.3814669>

- Objective dynamics: different objective combinations possess unique inherent balances. For instance, the trade-off between MSE and rule volume behaves differently than the relationship between MSE and information gain.
- Novelty as a driver: integrating a novelty score, specifically one derived from match sets, proves to be the most promising strategy for facilitating the generation of low-error, low-complexity models.

2 Related Work

Early explorations by Ishibushi et al. [19] used a multi-objective genetic algorithm to balance classification accuracy against model complexity, demonstrating that treating these as distinct objectives outperforms single-objective methods. Similarly, Bernadó i Mansilla and Garrell i Guíu [3] compare Pareto-based and non-Pareto-based rule evaluation strategies in the MOLeCS system. Their work confirmed that while multi-objective optimization lead to better results compared to accuracy-only approaches, a weighted balance that favors accuracy yielded the most robust results.

MO-XCS [8] introduced multi-objective reward vectors and policies to the XCS classifier system [33], while Djartov et al. [12] implemented a Michigan-style LCS using NSGA-II's non-dominated sorting to handle conflicting objectives in robust aviation decision making.

A recent architecture, the Heuristic Evolutionary Rule Optimization System (HEROS) [21], incorporates the Pareto-Inspired Multi-Objective Rule Fitness (PIMORF) [32]. Originally developed for use in ExSTraCS [29], PIMORF dynamically balances accuracy and coverage using the current objective space Pareto-front without prior knowledge of the problem noise. HEROS was recently expanded [20] to include a phase alternation scheme somewhat similar to SupRB's to facilitate anytime capabilities. Additionally, Lipschutz-Villa et al. explored the usage of tree-based initialization to give the rule discovery methods a better informed starting point. Overall, they showed that their extensions of HEROS can improve its performance on at least one axis of the two-objective problem.

Still, a significant research gap remains in using Pareto-optimality concepts for discovering effective rules efficiently. Most existing approaches still rely on mechanisms to implicitly achieve accurate but general rules or on multi-objective yet scalarized goals or even manual fitness weighting at some stage of evaluation. This research addresses that limitation by replacing traditional fitness functions entirely with a Pareto-based evaluation method. The proposed system enables a more dynamic exploration of the under-represented regions of the search space without the inherent constraints of manual weighting.

3 SupRB and Rule Discovery

SupRB [18] consists of two alternating phases: the RD phase in which newly generated individual rules are added to the global ever-expanding rule pool¹ and the SC phase where an optimal subset of the rule pool is constructed, forming a complete solution to the regression task. This cycle repeats until either an iteration limit or a termination criterion is met (cf. Figure 1).

¹After a rule has been added to the pool it remains unchanged throughout the remainder of the training process.

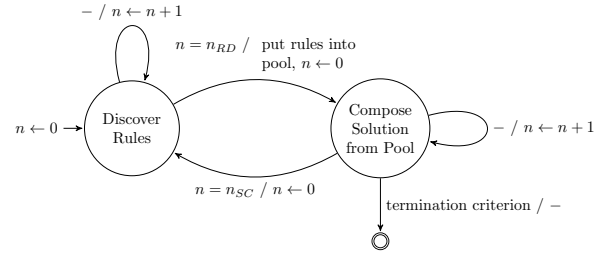


Figure 1: SupRB alternates Rule Discovery (RD) and Solution Composition (SC) phases until a termination criterion is met. RD generates new rules and adds them to the global rule pool. SC selects a minimal well-performing subset of rules from the rule pool and constructs a solution to the regression task. Adapted from Heider et al. [18].

By default and in most of the experiments so far [13], SupRB's rules consist of a hyperrectangular matching function and a linear regression model. The original RD uses a $(1, \lambda)$ -Evolution Strategy (ES) to discover new rules, generating a single rule at a time and independent of other rules². The initial individual of the ES is placed so that it exactly matches a randomly selected training example, i.e. lower and upper bounds are the same as the data point's values, prioritizing regions where the current elitist solution, i.e. the individual produced by the last SC phase as a subset of the rule pool, exhibits high error. From this rule, the ES generates λ children by expanding lower and upper bounds using normal mutation. In this version, the fitness of a child is a scalarized function which balances accuracy (measured in MSE) against the volume of the input space the child covers. Then, the best-performing child replaces the parent³ for the subsequent iteration. The process terminates if the elitist rule remains unimproved for δ iterations or if the set iteration limit is reached. This instantiation of RD serves as our baseline (denoted in later sections as BL).

Following the RD phase, the SC phase uses a Genetic Algorithm (GA) to assemble the individual rules in the rule pool into solution candidates⁴. Here, a candidate is defined as a specific subset of the rule pool. The GA evaluates these subsets based on two primary criteria: their in-sample MSE, and their complexity, defined by the number of rules selected.

4 Multi-objective Rule Discovery

In our approach, two objectives are optimized simultaneously to diversify rule generation and avoid the limitations of scalarization. This multi-objective framework allows for a dynamic balance between competing objectives, such as predictive performance and rule generality, without enforcing a fixed trade-off. The primary

²Because rules are discovered individually, this RD approach cannot be categorized as Michigan-style.

³If a parent that is replaced has a higher fitness than the new candidate, it gets stored in an archive and the best overall performing rule is returned at the end of the search.

⁴Rules remain unmodified by the GA and thus it can operate on a bit string with a bit for each rule in the pool. Consequently, this SC phase is not a typical Pittsburgh-style approach.

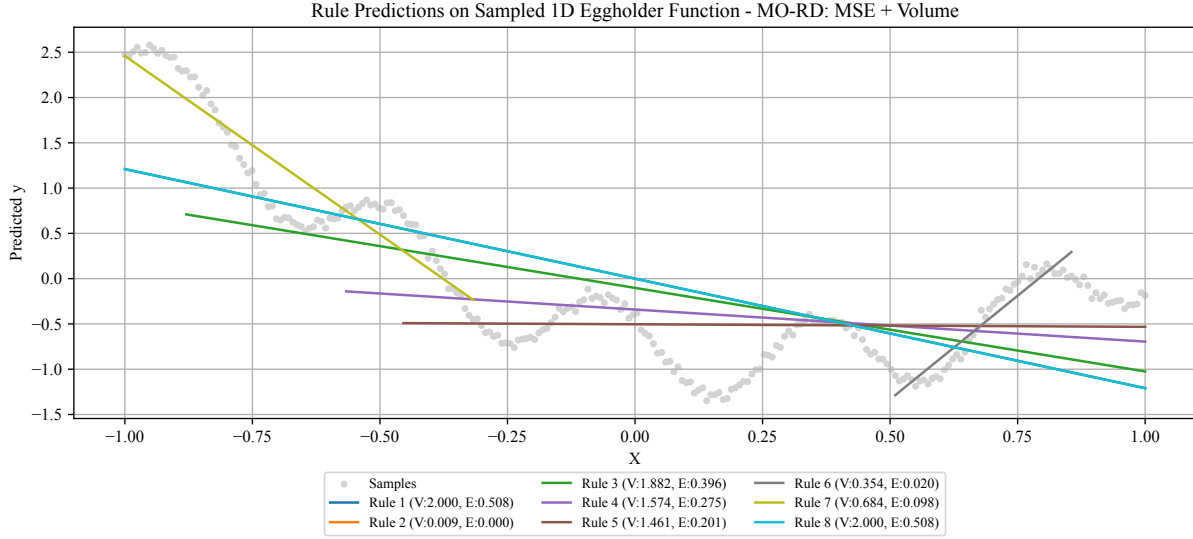


Figure 2: Function space coverage of the non-dominated set of rules generated on a one-dimensional Eggholder [10, 26] benchmark function with an MO-RD using volume (V) and MSE (E) as objectives. This shows that a large share of the Pareto set will automatically be non-effective for solving the learning task because having a high volume or unique tradeoff between volume and MSE is not necessarily useful. Interestingly, in later end-to-end experiments (cf. Figure 3), the generality metrics perform better than expected from the one-dimensional results which is likely caused by sparsity in higher dimensional spaces.

objective we chose is MSE, as low error is crucial for accurate predictions. The second objective can be chosen based on specific requirements, such as rule generality, complexity, or novelty, allowing flexibility in optimizing different aspects of rule quality. We replace the original ES with NSGA-II due to its simplicity and strong performance in two-objective optimization, comparable to other algorithms like SPEA-II [36] and MOEA-D [35]. However, unlike decomposition-based approaches, NSGA-II does not require prior knowledge of weights between objectives, making it more versatile for our scenario. Instead of parallelizing n_{rules} individual ESs, we evolve a single population of size μ , where $n_{\text{rules}} \leq \mu$, enabling efficient exploration of the objective space. To start the rule discovery process, the initial population is generated. Instead of a single origin rule like in the ES, a set of μ origin rules is generated. These rules are placed near data samples where the current elitist solution⁵ exhibits the highest in-sample error. In each iteration, λ parents are selected to generate one child each via mutation of their condition bounds. These children are then merged with the current population. The combined population is sorted using non-dominated sorting into domination fronts, where rules in the same front are non-dominated with respect to the two objectives. Crowding distance is calculated for rules within each front to maintain diversity and prevent premature convergence. The next generation is formed by adding entire fronts, starting with the non-dominated set $\mathcal{F}_0$. If a front exceeds the remaining population capacity, it is truncated by prioritizing rules with the largest crowding distance

⁵When initializing a SupRB run, a dummy elitist solution is first instantiated. This solution is defined by a null genome (e.g. [0,0,0,...]), resulting in maximum error across all samples. Consequently, during the first RD phase, the initial population is placed around random samples.

until the population size μ is reached. After a set number of iterations, the algorithm returns the first n_{rules} from the final population's non-dominated set. This approach ensures a diverse set of high-quality rules, balancing predictive performance and interpretability, which can then be used effectively in the SC phase to construct the final model.

5 Observations on Objective Trade-offs

We begin by implementing MO-RD using rule volume (V) as the secondary objective, mirroring the baseline ES approach. In the scalarized baseline, the objective function naturally balances accuracy and generality through an implicit trade-off: as a rule's condition bounds expand, the increase in error eventually outweighs the gain in volume, leading to a decrease in overall fitness⁶. In the multi-objective framework, this balancing mechanism is absent. Since accuracy and rule volume are treated as independent objectives, extremely general rules with high error remain Pareto-optimal. Within the objective space, these rules are not dominated by more precise rules; instead, they occupy a different extreme of the trade-off. As shown in Figure 2, this initially suggests a significant generality problem, where the rule pool becomes populated with overly broad rules that appear too imprecise for effective SC.

We also evaluate rule support (S) as an alternative to rule volume. Rule support is defined as the number of data samples contained within the rule's volume, incorporating the underlying sample density. A rule can maintain high support if it covers a dense region, while a general rule may exhibit low support if it spans

⁶A key component of applying this strategy is to start with small rules and then only expand them which ensures that the originality targeted training example (region with a weaker prediction) remains covered

a sparse region. Despite this nuance, rule support exhibits similar characteristics to rule volume. Without explicit penalties for error or redundancy, the search produces high-support, high-error generalists that dominate the Pareto front.

Interestingly, our end-to-end experiments (cf. Section 5.3) reveal that the presence of these overly general rules is less detrimental than the isolated experiments in Figure 2 suggest. While these rules dominate large portions of the search space, the SC phase acts as a final filter, effectively ignoring high-error generalists when more specialized, accurate rules are available (as seen in Figures 3a and 3b, where the naive approach is denoted as V and S). This suggests that the MO-RD approach where the scalarized fitness functions is directly translated to Pareto-front approximation is more robust than expected. We hypothesize that the system's final performance is less dictated by the average quality of the rule pool but primarily by its diversity.

To test this hypothesis and address the structural limitations of generality-based metrics, we shift the focus of the RD process toward generating a rule pool that balances accuracy with diversity. This approach prioritizes the creation of rules that not only minimize error but also explore distinct regions of the input space, ensuring a more comprehensive and flexible set of rules for the SC phase.

5.1 Diversity Objectives

To evolve a comprehensive rule set, we define objectives targeting two complementary aspects: target homogeneity and population diversity. Information gain and variance reduction serve as implicit diversity promoters by measuring how effectively a rule partitions the data distribution to isolate consistent target values. These objectives may encourage exploration of distinct statistical patterns across the input space. Both objectives were inspired by their use in decision trees [24]. While we expect that their performance in decision trees does not translate to our usecase, we still expect them to provide some pressure into a useful direction.

We also explore the role of explicit diversity maintenance by incorporating a novelty score into MO-RD. Previously introduced by Heider et al. [16, 17] for novelty search in SupRB, this score was designed specifically for quantifying dissimilarity between rules. While novelty search performed similar to ES baseline in their study, we integrate it here to test its efficacy (and possible superiority) within our MO-RD approach.

All objectives are formulated as minimization problems; since higher information gain, variance reduction, and novelty values indicate superior rule quality, we minimize their negatives.

The information gain (IG) fitness objective is defined as

$$f_{IG}(m) = -(H - [pH_m + (1 - p)H_{\neg m}]), \quad (1)$$

where m is the binary matching mask, representing rule coverage across the dataset, H is the initial Shannon entropy of the dataset target values, and H_m is the Shannon entropy of target values for samples matched by a specific rule. The term p denotes the rule coverage, representing the proportion of samples covered by the rule

$$p = \frac{1}{K} \sum_{i=1}^K m_i, \quad (2)$$

where K is the total number of samples and $m_i \in \{0, 1\}$ is the i -th element of the binary matching mask, indicating whether the i -th sample is located inside the rule's condition bounds.

The entropy for a specific partition $S \in \{m, \neg m\}$ is calculated as

$$H_S = - \sum_{j=1}^C p_j \log_2(p_j + \epsilon), \quad (3)$$

where C is the number of unique target values present in the subset, p_j is the relative frequency of the j -th unique value, and ϵ (10^{-12}) is a small constant used for numerical stability.

While information gain is traditionally a classification metric, it is employed here as a strict measure of target value consistency. In the context of a regression task, $f_{IG}(m)$ treats unique target values as discrete states, allowing the rule discovery process to isolate regularities where the target remains constant or highly repetitive.

We also implement a variance reduction (VR) objective as a direct alternative to information gain, as it is more conventionally suited for regression tasks. This objective measures the reduction in target variance achieved by partitioning the data with a specific rule, rewarding rules that isolate samples with similar numerical values:

$$f_{VR}(m) = -(\sigma^2 - [p\sigma_m^2 + (1 - p)\sigma_{\neg m}^2]), \quad (4)$$

where σ^2 is the global variance of the target variable y , and σ_m^2 and $\sigma_{\neg m}^2$ are the variance of the target values for the samples matched or not matched by the specific rule, respectively. The calculation of p and p_j is analogous to the $f_{IG}(m)$ objective. The variance of any subset S is calculated as

$$\sigma_m^2 = \frac{1}{|S|} \sum_{y_i \in S} (y_i - \bar{y}_S)^2, \quad (5)$$

where $\bar{y}_S$ is the mean of the target values within that subset. While the information gain objective acts as a strict filter for target consistency, the variance reduction objective offers a smoother gradient, rewarding rules that find regions where target values are similar.

The novelty score (N) directly quantifies how dissimilar a rule's matching mask, i.e. which training examples are matched by the rule represented as a bit array, is compared to its neighbouring peers in the current population $\mathcal{A}$.⁷

$$f_N(u) = - \left(\frac{1}{k} \sum_{v \in \mathcal{N}_k(u, \mathcal{A})} d(u, v) \right). \quad (6)$$

It is defined as the average Hamming distance d between rule u and its k -nearest neighbors v within the population $\mathcal{A}$ [17]. In the discovery process, the novelty score encourages the evolution of rules that cover different regions of the input space.

5.2 Experimental Setup

Two-phase end-to-end experiments are conducted to evaluate each fitness objective combination, where each combination consists of MSE plus either a generality or diversity measure. The general

⁷We tested an extended version of the archive $\mathcal{A}$ that included rules from previous generations. However, this approach led to excessive memory usage when the archive exceeded 256 rules. Our experiments suggest that for the extended version to yield significant improvements in rule diversity, its size would need to be increased substantially.

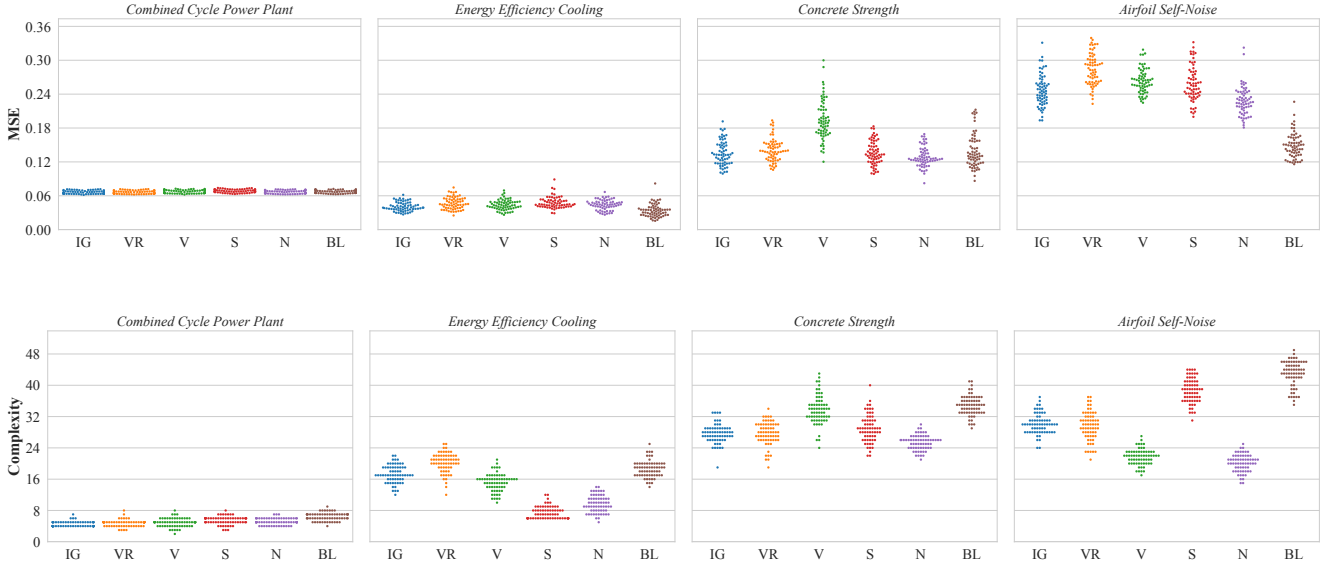


Figure 3: Swarm plots of final model in-sample MSEs (top) and rule count within the model (complexity; bottom). Ordered from easiest (left) to most difficult (right) dataset to be optimized for. The configurations’ results especially differ on the more difficult datasets as only these allow the advanced metrics to unfold their capabilities.

Configurations shown: information gain (IG), variance reduction (VR), volume (V), support (S), novelty score (N), and the baseline (BL).

experimental setup and dataset choice are heavily based on Heider [13] to increase comparability to these results. Specifically, the Airfoil Self-Noise (ASN) [4], Energy Efficient Cooling (EEC) [27], Combined Cycle Power Plant (CCPP) [28], and Concrete Strength (CS) [34] datasets are used to evaluate each SupRB setup on real-world linear and non-linear regression problems.

The experiments are conducted in two phases. In the first phase, we use the Optuna [1] for hyperparameter tuning. While the rule fitness evaluation is based on Pareto-dominance, we still tune the population size μ , number of children generated in each iteration λ , and the RD iteration limit n_{iter} itself. The solution composition’s hyperparameters remain fixed. A limit of a 1000 Optuna trials per dataset is set, where each trial consists of a different hyperparameter configuration. Using four-fold cross validation, each trial is evaluated on its final model’s in-sample MSE.

In the second experiment phase, the best-performing hyperparameter configurations of the first phase are reevaluated on each dataset using 8-fold Monte-Carlo cross validation⁸. These 64 runs per configuration are the basis for the following *Results* section and its plots.

As a baseline, SupRB with its standard ES rule discovery is tuned and evaluated in the same way. The hyperparameters tuned include λ , n_{iter} , and the fitness weight α .

Both the experiment and project codebases⁹ are provided as supplementals.

5.3 Results

The experimental evaluation across the CCPP, EEC, CS, and ASN datasets demonstrates that MO-RD provides a competitive alternative to the baseline (BL). As illustrated in Figure 3, the performance profiles of the various configurations diverge as the complexity of the underlying data distribution increases, with CCPP as the simplest dataset to optimize for, and ASN as the most difficult.

On the CCPP, EEC, and CS datasets, MO-RD configurations achieve MSEs comparable to the baseline while maintaining lower model complexity. However, on the ASN dataset, the baseline configuration achieves a lower MSE than any MO-RD variant. This higher predictive performance is accompanied by a substantial increase in model complexity. The baseline requires approximately double the amount of rules (typically 40 – 48) compared to the novelty (N), and volume (V) configurations (typically 18 – 24). This suggests that optimizing solely for MSE in the hyperparameter tuning phase of the experiments encourages the baseline to discover mainly high-accuracy low-volume rules, whereas MO-RD variants prioritize more generalized structures that allow SC to select a rule set that is easier to interpret.

⁸The composition of individual datasplits did not have any influence on the results, i.e. error distributions are similar across splits.

⁹Both codebases have been anonymized with the git logs removed but will be released publicly on Github if the paper is accepted.

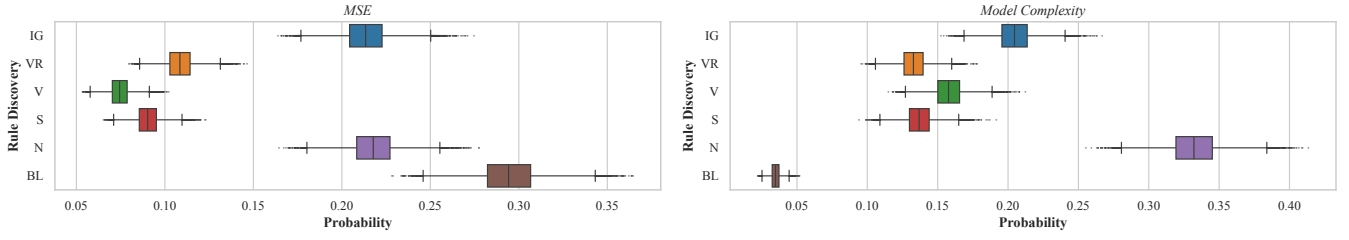


Figure 4: Box plots of posterior distributions (for a given run achieving the best performance among the group of tested methods) of model MSE and complexities respectively. Model adapted from Calvo et al. [6, 7]. 1.5 IQR and outliers are shown. For the MSE, the baseline (BL), SupRB with novelty search RD (N), or information gain RD (IG) are the most likely candidates to produce the best model (on similar datasets as we investigated), whereas for complexity, the baseline is not a good candidate but novelty search RD seems to be the best choice. Intuitively, novelty search RD seems to produce rule pools that feature less overlap but can be combined to similar system performance.

The novelty score objective demonstrated the most consistent performance across all datasets. By rewarding rules located in underrepresented regions of the input space, it increases overall diversity in the RD rule population and later in the global rule pool. The effectiveness of the novelty score configuration directly aligns with prior research by Heider et al. [17, 13], who observed that, in single objective RD, what they called NSLC-P (novelty search with local competition, i.e. still using the scalarized fitness for local decisions of similarly novel rules and where novelty is measured against the pool rather than the current population) and MCNS (minimal criteria novelty search, i.e. at least some prediction power and experience is required for a rule to be competitive) achieved comparable performance to the evolution strategy used in the baseline SupRB. These results suggest that novelty-based fitness drives population diversity and enhances parsimony and might be able to enhance systems beyond SupRB as well.

Information gain (IG) and variance reduction (VR) exhibit degrading performance with increasing dataset difficulty. While effective for local refinement, the results suggest that their greedy nature often leads to stagnation in local optima. On non-linear datasets, like CS and ASN, they are consistently outperformed by the other MO-RD variants.

Although initial isolated experiments suggested that volume (V) and support (S) might prioritize the generation of overly broad rules (cf. Section 5), the full SupRB runs indicate that these objectives remain viable. The MO-RD appears to generate a sufficiently large set of high-accuracy rules that enables the construction of accurate models, even when an objective directly favors generality.

We employ the Bayesian Plackett-Luce ranking model proposed by Calvo et al. [6, 7] for algorithm performance analysis for the statistical evaluation of our experiment results. Given rankings over k configurations induced by some metric for each dataset and random seed, the model supplies a posterior distribution over its k parameters, which can directly be interpreted as the probability that the given configuration performs overall best in terms of the metric (i.e., the win probability). This deliberately ignores the magnitude of differences, and therefore allows pooling of factors where the metric itself is not comparable (e.g., MSE between different datasets). We perform this analysis twice, using MSE and model

complexity as metric, respectively, with lower values achieving a higher ranking.

Similarly to experiments in the work of Heider [13], no configuration reached a sufficiently high win probability to justify a clear selection as the best performing one. Still, some observations can be made. Although variance reduction and information gain show comparable performance in Figure 3, the analysis shows that among the investigated group of objectives and datasets, information gain maintains a higher win probability of the two. The baseline is a good candidate to achieve the best overall MSE, primarily informed by its strong performance on the ASN dataset, while the novelty score configuration maintains the highest probability to minimize model complexity. Based on Figure 4, the novelty score configuration appears to be the most promising candidate for MO-RD based on its largest win probabilities for model complexity and its competitive win probabilities for MSE.

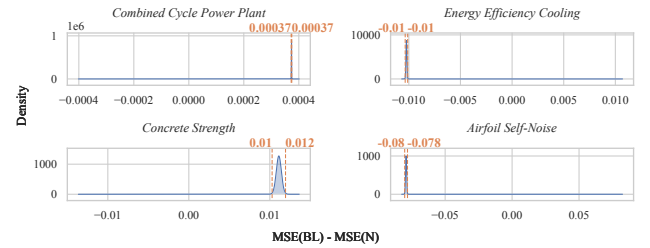


Figure 5: Density plot of the posterior distribution of the Bayesian correlated t-test derived from Corani and Benavoli [9] for model MSEs between N and BL configurations. Previously used by Heider [13] to evaluate the differences in SupRB configurations. The orange dashed lines denote the 99% Highest Posterior Density Interval (HPDI), meaning that 99% of the posterior probability mass lie within these bounds. In essence, the plot shows the likely range of differences between single runs. A probability mass above zero indicates that the BL configuration yields a higher MSE than the N configuration. The results are rounded to two significant figures.

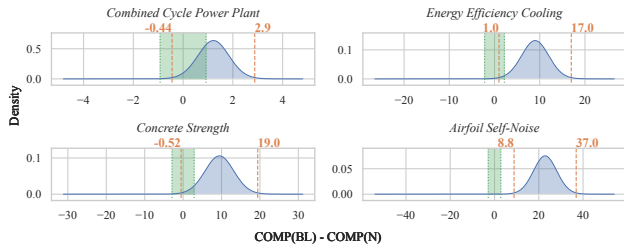


Figure 6: Plot of the correlated t-test for test model complexities of the baseline (BL) and novelty score (N) configurations, analogous to Figure 5. The Region of Practical Equivalence (ROPE) shows where the configurations are expected to have equivalent performance. The region is highlighted in green. Across all datasets, the novelty score configurations is expected to construct smaller models.

To detail the effectiveness of MO-RD using the novelty score, we apply Corani and Benavoli’s correlated t-tests [9]. Here, we compare the novelty score and baseline configurations, allowing for a more precise probabilistic assessment of performance differences and effect size estimates on each dataset.

The t-test in Figure 6 shows a definitive difference in expected model complexity. Included is a Region of Practical Equivalence (ROPE)¹⁰, shown in green, highlighting the probability mass where the performance difference between configurations is negligible. On all datasets, the majority of the probability mass is located on the right of the ROPE, meaning that the novelty score configuration is expected to produce smaller models than the baseline.

Figure 5 shows the correlated t-test for in-sample MSE. Across most datasets, the expected performance differences are marginal, albeit not clearly one-sided. For the CCPP, EEC, and CS datasets, this difference remains below 0.012 standard deviations. The most pronounced discrepancy occurs on the ASN dataset, where the difference reaches approximately 0.08 of a standard deviation which is statistically significant but might still not matter in practical application. While these results suggest near equivalence we refrain from establishing a ROPE for MSE as there is no well-established way to assess this in an application-independent way. In high-precision ML applications, even these nominal variations in error might be operationally significant which has to be decided based on stakeholder input. Conversely, defining a ROPE is appropriate for model complexity, as small increases in complexity typically do not result in a detrimental loss of interpretability.

The experimental results observed in SupRB demonstrate that a novelty-driven diversity metric can significantly improve rule discovery. However, the portability of this mechanism depends on the underlying structure of the RBML architecture used. While SupRB’s RD phase evaluates rules individually, other frameworks present challenges and opportunities for implementing similar novelty scores.

In systems like GAssist [2], a novelty score cannot be easily implemented due to its Pittsburgh-style nature. GAssist evaluates the fitness of entire rule sets instead of individual rules. Therefore, novelty scores that evaluate dissimilarity between individual rules are

structurally mismatched with GAssist’s holistic evaluation. Integrating novelty here would require a fundamental shift in how the system decomposes set-level performance into individual rule contributions.

Michigan-style architectures, especially those derived from XCS where multi-objective evaluation was already tested (e.g. MO-XCS), can use novelty scores for the preservation of rules that cover underrepresented regions of the input space. By rewarding rules that cover niche spaces, the system ensures that the deletion scheme does not prune these specialized rules simply because they have not been used recently. This ensures that potentially high-value rules are preserved and might protect against cover-delete-cycles.

6 Conclusion and Future Work

As AI becomes increasingly integrated into critical systems, the demand for models that are simultaneously high-performing and human-interpretable has never been higher. RBML provides an explainable alternative to black-box systems, yet its efficacy often hinges on finding the right balance between predictive performance and the model’s rule count, i.e. its complexity. In this paper, we addressed the challenge of providing good rules for model selection by introducing a Multi-Objective Rule Discovery (MO-RD). We implemented this approach in the modular SupRB framework to show the efficacy in a recent RBML system.

By replacing SupRB’s original rule discovery with NSGA-II, we transitioned from a scalar fitness function to a Pareto-based method. This shift allows for the simultaneous optimization of accuracy (MSE) alongside secondary objectives like diversity or generality. Our evaluation across four real-world regression tasks, backed by Bayesian statistical analysis, demonstrates that MO-RD configurations, particularly those employing a novelty score based on rules’ match sets, achieve predictive performance compared to standard SupRB while significantly reducing model complexity.

This work highlights that Pareto-optimality can serve as the primary driver for the rule discovery process itself. By removing the reliance on manual, scalarized biases, the system explores the rule space more dynamically, facilitating the evolution of diverse and efficient rule sets.

The findings suggest that the definition of fitness objectives may be a more significant factor in system success than the specific optimization metaheuristic chosen. Consequently, future work will focus on:

- Alternative utility metrics: exploring new indicators of rule utility to further refine rule evaluation.
- Integrated feedback loops: investigating the interplay between rule discovery and solution composition to determine how current model deficiencies can autonomously guide the evolution of rules toward specific gaps in the solution space.
- Architecture portability: implementing the multi-objective rule discovery approach within other RBML architectures to validate its general applicability.

Ultimately, this work establishes an approach for maintaining predictive performance while improving model interpretability by reducing model size, addressing a core requirement for the practical application of RBML.

¹⁰More details in Heider [13].

References

- [1] Takuya Akiba, Shotaro Sano, Toshihiko Yanase, Takeru Ohta, and Masanori Koyama. 2019. Optuna: a next-generation hyperparameter optimization framework. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining (KDD '19)*. Association for Computing Machinery, Anchorage, AK, USA, 2623–2631. ISBN: 9781450362016. doi: 10.1145/3292500.3330701.
- [2] Jaume Bacardit Peñarroya. 2004. *Pittsburgh genetic-based machine learning in the data mining era: representations, generalization, and run-time*. Ph.D. Dissertation. Universitat Ramon Llull.
- [3] Ester Bernadó i Mansilla and Josep M. Garrell i Guíu. 2001. MOLECS: using multiobjective evolutionary algorithms for learning. In *Evolutionary Multi-Criterion Optimization*. Eckart Zitzler, Lothar Thiele, Kalyanmoy Deb, Carlos Artemio Coello Coello, and David Corne, (Eds.) Springer Berlin Heidelberg, Berlin, Heidelberg, 696–710. ISBN: 978-3-540-44719-1.
- [4] Thomas F. Brooks, D. Stuart Pope, and Michael A. Marcolini. 1989. Airfoil Self-Noise and Prediction. Techreport 1218.
- [5] Martin V. Butz, Tim Kovacs, Pier Luca Lanzi, and Stewart W. Wilson. 2001. How XCS evolves accurate classifiers. In *Proceedings of the 3rd Annual Conference on Genetic and Evolutionary Computation (GECCO'01)*. Morgan Kaufmann Publishers Inc., San Francisco, California, 927–934. ISBN: 1558607749.
- [6] Borja Calvo, Josu Ceberio, and Jose A. Lozano. 2018. Bayesian inference for algorithm ranking analysis. In *Proceedings of the Genetic and Evolutionary Computation Conference Companion (GECCO '18)*. Association for Computing Machinery, New York, NY, USA, 324–325. ISBN: 978-1-4503-5764-7. doi: 10.1145/3205651.3205658.
- [7] Borja Calvo, Ofer M. Shir, Josu Ceberio, Carola Doerr, Hao Wang, Thomas Bäck, and Jose A. Lozano. 2019. Bayesian performance analysis for black-box optimization benchmarking. In *Proceedings of the Genetic and Evolutionary Computation Conference Companion (GECCO '19)*. Association for Computing Machinery, New York, NY, USA, 1789–1797. ISBN: 978-1-4503-6748-6. doi: 10.1145/3319619.3326888.
- [8] Xiu Cheng, Gang Chen, and Mengjie Zhang. 2016. An xcs-based algorithm for multi-objective reinforcement learning. In *2016 IEEE Congress on Evolutionary Computation (CEC)*. doi: 10.1109/CEC.2016.7744298.
- [9] Giorgio Corani and Alessio Benavoli. 2015. A Bayesian approach for comparing cross-validated algorithms on multiple data sets. *Machine Learning*, 100, 2, 285–304. doi: 10.1007/s10994-015-5486-z.
- [10] Jacek M. Czerwik and Hubert Zarzycki. 2017. Artificial acari optimization as a new strategy for global optimization of multimodal functions. *Journal of Computational Science*, 22, 209–227. doi: 10.1016/j.jocs.2017.05.028.
- [11] K. Deb, A. Pratap, S. Agarwal, and T. Meyarivan. 2002. A fast and elitist multi-objective genetic algorithm: NSGA-II. *IEEE Transactions on Evolutionary Computation*, 6, 2, 182–197. doi: 10.1109/4235.996017.
- [12] Boris Djartov and Sanaz Mostaghim. 2023. Multi-objective Multiplexer Decision Making Benchmark Problem. In *Proceedings of the Companion Conference on Genetic and Evolutionary Computation (GECCO '23 Companion)*. Association for Computing Machinery, New York, NY, USA, 1676–1683. ISBN: 979-8-4007-0120-7. doi: 10.1145/3583133.3596360.
- [13] Michael Heider. 2025. *Disentangling the Model Selection Tasks for Improved Explainability in a Rule-Based Machine Learning System*. Ph.D. Dissertation. Universität Augsburg.
- [14] Michael Heider, David Pätz, Helena Stegherr, and Jörg Hähner. 2023. A Metaheuristic Perspective on Learning Classifier Systems. In *Metaheuristics for Machine Learning*. Mansour Eddaly, Bassem Jarboui, and Patrick Siarry, (Eds.) Springer Nature Singapore, Singapore, 73–98. ISBN: 978-981-19-3888-7. doi: 10.1007/978-981-19-3888-7_3.
- [15] Michael Heider, David Pätz, and Alexander R. M. Wagner. 2022. An overview of LCS research from 2021 to 2022. In *Proceedings of the Genetic and Evolutionary Computation Conference Companion*, 2086–2094. doi: 10.1145/3520304.3533985.
- [16] Michael Heider, Helena Stegherr, David Pätz, Roman Sraj, Jonathan Wurth, Benedikt Volger, and Jörg Hähner. 2022. Approaches for Rule Discovery in a Learning Classifier System: in *Proceedings of the 14th International Joint Conference on Computational Intelligence*. SCITEPRESS - Science and Technology Publications, Valtella, Malta, 39–49. ISBN: 978-989-758-611-8. doi: 10.5220/0011542000003332.
- [17] Michael Heider, Helena Stegherr, David Pätz, Roman Sraj, Jonathan Wurth, Benedikt Volger, and Jörg Hähner. 2023. Discovering Rules for Rule-Based Machine Learning with the Help of Novelty Search. *SN Computer Science*, 4, 6, 778. doi: 10.1007/s42979-023-02198-x.
- [18] Michael Heider, Helena Stegherr, Jonathan Wurth, Roman Sraj, and Jörg Hähner. 2022. Separating rule discovery and global solution composition in a learning classifier system. In *Proceedings of the Genetic and Evolutionary Computation Conference Companion*. ACM, 248–251. ISBN: 978-1-4503-9268-6. doi: 10.1145/3520304.3529014.
- [19] Hisao Ishibuchi, Tadahiko Murata, and I.B. Türkşen. 1997. Single-objective and two-objective genetic algorithms for selecting linguistic rules for pattern classification problems. *Fuzzy Sets and Systems*, 89, 2, 135–150. doi: 10.1016/S0165-0114(96)00098-X.
- [20] Gabriel Lipschutz-Villa, Harsh Bandhey, Khoi Dinh, Michael Heider, Malek Kamoun, and Ryan J. Urbanowicz. 2026. Phase-alternation and tree-initialization to facilitate interpretable rule-based machine learning. In *Proceedings of the Genetic and Evolutionary Computation Conference*. in press.
- [21] Gabriel Lipschutz-Villa, Harsh Bandhey, Ruonan Yin, Malek Kamoun, and Ryan Urbanowicz. 2025. Rule-based Machine Learning: Separating Rule and Rule-Set Pareto-Optimization for Interpretable Noise-Agnostic Modeling. In *Proceedings of the Genetic and Evolutionary Computation Conference*. ACM, NH Malaga Hotel Malaga Spain, 407–415. ISBN: 979-8-4007-1465-8. doi: 10.1145/3712256.3726461.
- [22] Yi Liu, Will N. Browne, and Bing Xue. 2021. A comparison of learning classifier systems' rule compaction algorithms for knowledge visualization. *ACM Trans. Evol. Learn. Optim.*, 1, 3, Article 10, 38 pages. doi: 10.1145/3468166.
- [23] David Pätz, Richard Nordsieck, and Jörg Hähner. 2024. Measuring similarities in model structure of metaheuristic rule set learners. In *Applications of Evolutionary Computation*. Stephen Smith, João Correia, and Christian Cintrano, (Eds.) Springer Nature Switzerland, Cham, 256–272. ISBN: 978-3-031-56855-8.
- [24] Alexey Poyarkov, Alexey Drutsa, Andrey Khalyavin, Gleb Gusev, and Pavel Serdyukov. 2016. Boosted decision tree regression adjustment for variance reduction in online controlled experiments. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '16)*. Association for Computing Machinery, San Francisco, California, USA, 235–244. ISBN: 9781450342322. doi: 10.1145/2939672.2939688.
- [25] Abubakar Siddique, Michael Heider, Muhammad Iqbal, and Hiroki Shiraishi. 2024. A Survey on Learning Classifier Systems from 2022 to 2024. In *Proceedings of the Genetic and Evolutionary Computation Conference Companion*. Association for Computing Machinery, 1797–1806. ISBN: 979-8-4007-0495-6. doi: 10.1145/3638530.3664165.
- [26] Anthony Stein, Simon Menssen, and Jörg Hähner. 2018. What about interpolation? a radial basis function approach to classifier prediction modeling in XCSF. In *Proceedings of the Genetic and Evolutionary Computation Conference (GECCO '18)*. Association for Computing Machinery, Kyoto, Japan, 537–544. ISBN: 9781450356183. doi: 10.1145/3205455.3205599.
- [27] Athanasios Tsanas and Angeliki Xifara. 2012. Accurate quantitative estimation of energy performance of residential buildings using statistical machine learning tools. *Energy and Buildings*, 49, 560–567. doi: 10.1016/j.enbuild.2012.03.003.
- [28] Pınar Tüfekci. 2014. Prediction of full load electrical power output of a base load operated combined cycle power plant using machine learning methods. *International Journal of Electrical Power & Energy Systems*, 60, 126–140. doi: 10.1016/j.ijepes.2014.02.027.
- [29] Ryan Urbanowicz and Jason Moore. 2015. Retooling fitness for noisy problems in a supervised michigan-style learning classifier system. In *Proceedings of the 2015 Annual Conference on Genetic and Evolutionary Computation (GECCO '15)*. Association for Computing Machinery, Madrid, Spain, 591–598. ISBN: 9781450334723. doi: 10.1145/2739480.2754756.
- [30] Ryan J. Urbanowicz and Will N. Browne. 2017. *Introduction to Learning Classifier Systems*. SpringerBriefs in Intelligent Systems. Springer Berlin Heidelberg, Berlin, Heidelberg. ISBN: 978-3-662-55006-9 978-3-662-55007-6. doi: 10.1007/978-3-662-55007-6.
- [31] Ryan J. Urbanowicz and Jason H. Moore. 2009. Learning Classifier Systems: A Complete Introduction, Review, and Roadmap. *Journal of Artificial Evolution and Applications*.
- [32] Ryan J. Urbanowicz, Randal S. Olson, and Jason H. Moore. 2016. Pareto inspired multi-objective rule fitness for noise-adaptive rule-based machine learning. In *Parallel Problem Solving from Nature – PPSN XIV*. Julia Handl, Emma Hart, Peter R. Lewis, Manuel López-Ibáñez, Gabriela Ochoa, and Ben Paechter, (Eds.) Springer International Publishing, 514–524. ISBN: 978-3-319-45823-6.
- [33] Stewart W. Wilson. 1995. Classifier Fitness Based on Accuracy. *Evolutionary Computation*, 3, 2, 149–175. doi: 10.1162/evco.1995.3.2.149.
- [34] I.-C. Yeh. 1998. Modeling of strength of high-performance concrete using artificial neural networks. *Cement and Concrete Research*, 28, 12, 1797–1808. doi: 10.1016/S0008-8846(98)00165-3.
- [35] Qingfu Zhang and Hui Li. 2007. MOEA/D: a multiobjective evolutionary algorithm based on decomposition. *IEEE Transactions on Evolutionary Computation*, 11, 6, 712–731. doi: 10.1109/TEVC.2007.892759.
- [36] Eckart Zitzler, Marco Laumanns, and Lothar Thiele. 2001. Spea2: improving the strength pareto evolutionary algorithm. *TIK report*, 103.

Received 3 April 2026; revised 5 May 2026; accepted 26 April 2026