

# Decision Rule-Based Learning of Terrorist Threats

Nida Meddouri<sup>1</sup><sup>a</sup>, Loïc Salmon<sup>2</sup><sup>b</sup>, David Beserra<sup>1</sup><sup>c</sup> and Elloh Adja<sup>1</sup>

<sup>1</sup>*Laboratoire de Recherche de l'EPITA, Le Kremlin-Bicêtre, France*

<sup>2</sup>*Institut des Sciences Exactes et Appliquées, University of New-Caledonia, France*  
fi - - - -

**Keywords:** Data Mining, Machine Learning, Decision Rule, Criminality, Terrorist Threats.

**Abstract:** Artificial Intelligence (AI) offers powerful tools for analyzing criminal data and predicting security threats. This paper focuses on the interpretable prediction of terrorist threats in France using official crime datasets from 2012 to 2021. We propose a preprocessing methodology to aggregate and label spatio-temporal crime data at the departmental level, addressing challenges such as data imbalance and structural heterogeneity. To ensure explainability, we adopt symbolic learning approaches based on decision rule generators implemented in WEKA, including MODLEM, NNge, and MOEFC. We evaluate these models through nine experiments simulating real-world prediction scenarios, using metrics such as misclassification rate, Recall, Kappa statistic, AUC-ROC, and AUPR. Results show that rule-based models achieve stable performance across periods, with Recall averaging 96% and AUPR close to 0.93, despite severe class imbalance. Among the tested methods, NNge and MOEFC provide the best trade-off between interpretability and predictive accuracy. These findings highlight the potential of interpretable rule-based models for supporting counter-terrorism strategies.

## 1 INTRODUCTION

Over the past two decades, criminal activity in France has evolved, leading to a significant rise in acts of malice, particularly in connection with social and labor movements, riots, and terrorism (Mucchielli, 2008). In this complex landscape, integrating artificial intelligence techniques presents promising opportunities to enhance public and private security systems. Studies conducted in various countries, including Brazil (Da Silva et al., 2020), the Middle East (Tolan et al., 2015), and others (Saidi and Trabelsi, 2022), have already demonstrated the effectiveness of spatio-temporal crime data analysis in this domain. Building on this foundation, this work aims to adapt and apply these approaches to the French context, focusing on developing an interpretable and explainable terrorism threat prediction model, leveraging a recent research (Meddouri and Beserra, 2024).

This work does not aim to introduce a new algorithm but rather to address a critical gap in the literature: the lack of interpretable and explainable models for predicting terrorist threats using real-world, highly imbalanced crime data. Our contribution lies

in (i) designing a reproducible preprocessing pipeline for spatio-temporal aggregation of official French crime datasets, (ii) systematically benchmarking a diverse set of state-of-the-art decision rule learners under severe class imbalance, and (iii) providing an interpretability-driven evaluation framework based on rule complexity and similarity analysis. To preserve interpretability, we deliberately avoided oversampling or synthetic data generation and instead relied on evaluation metrics robust to imbalance, such as Recall, AUPR, and Kappa statistic. These aspects are essential for operational decision-making in security contexts, where black-box models are often unsuitable.

In section 2, we present the record of criminality in France. In section 3, we present the analysis and preprocessing of criminality data and the challenge to discover. In section 4, we propose the interpretable learning of terrorist attacks in France. Finally, in section 5, we present an experimental study based on interpretable and explainable machine learning methods (rules generators) from the labeled criminality data.

<sup>a</sup> <https://orcid.org/0000-0002-7815-630X>

<sup>b</sup> <https://orcid.org/0000-0002-7267-6371>

<sup>c</sup> <https://orcid.org/0000-0002-7450-8531>

## 2 CRIME DATA IN FRANCE (2012-2021)

Since October 9, 2015, crime-related data in France has been available online under an Open Licence<sup>1</sup>. Covering the period from 2012 to 2021, these datasets encompass Metropolitan France, the *Overseas Departments and Regions*, and the *Overseas Collectivities*. They provide crime and offense statistics recorded by the national police and gendarmerie. Offenses are grouped into seven major categories: offenses against individuals, offenses against property, drug-related offenses, offenses against public authority, offenses related to public health and the environment, offenses under labor and competition law, and administrative and documentary offenses. This database is a valuable resource for spatio-temporal crime analysis, facilitating the modeling and interpretation of crime trends, particularly the emergence of specific offenses such as terrorist attacks.

The use and interpretation of these data require consideration of several key factors. First, the statistics account only for crimes and offenses, excluding minor infractions. These acts are recorded at the time they are first reported to security forces and brought to the attention of judicial authorities. Additionally, traffic offenses are not included in these counts.

Offenses are recorded by the administrative authority that observes and documents them. However, an offense is not necessarily reported or recorded in the same location where it was committed. This discrepancy particularly affects the *Public Security Districts*, the *Departmental Gendarmerie Units*, the *Border Police*, and judicial police services such as the *Central Directorate of the Judicial Police*, the *Regional Directorates of the Judicial Police*, and the *National Directorate of the Judicial Police*. This also applies to the *Republican Security Companies*, which operate across multiple departments or regions, as well as to central offices with national jurisdiction. Consequently, it is essential to distinguish between the number of offenses recorded by a given service and the actual number of crimes and offenses committed in the territory where that service is based.

Recorded offenses correspond to incidents documented within a given year. However, some offenses may have been committed in the previous year or, more rarely, even earlier but are accounted for in the year they were recorded. Conversely, offenses occurring late in the year may appear in the records of the following year.

<sup>1</sup><https://www.data.gouv.fr/fr/datasets/crimes-et-delits-enregistres-par-les-services-de-gendarmerie-et-de-police-depuis-2012/information>

Depending on the type of offense, these data may not fully reflect the level of insecurity perceived by citizens. For offenses without direct physical or moral victims—such as drug-related violations, labor law infractions, immigration offenses, environmental crimes, or prostitution-related offenses—the recorded figures primarily indicate law enforcement activity rather than the actual prevalence of such crimes. These numbers reflect the intensity of efforts to detect and prosecute offenses rather than a direct measure of criminal trends.

The unit of measurement varies by offense type, with each category assessed using the most relevant metric. However, this inconsistency in measurement methods makes direct aggregation of figures across different categories inappropriate.

Additionally, crime recording systems have undergone significant changes in recent years. Consequently, some variations in statistical trends result from modifications in data collection practices rather than actual shifts in criminal activity.

Lastly, the organization of gendarmerie and police services evolves over time, with jurisdictions being created, abolished, or restructured. These changes can complicate the interpretation of crime figures for a given service. Modifications to service jurisdictions are officially published in the *Journal Officiel*.

This database serves as a crucial resource for understanding and analyzing the evolution of crime in France over time. We believe that these data can be leveraged for spatio-temporal analyses to both “predict” and “interpret” the occurrence of terrorist attacks in France.

## 3 ANALYSIS AND PREPROCESSING OF CRIMINAL DATA

Before using these data for learning purposes, pre-processing is necessary to merge the statistics provided by the *National Police* and the *National Gendarmerie* on a year-by-year basis. The police services are organized into directorates (either national or specific to the Paris metropolitan area), each with its own territorial structure. In contrast, the organization of gendarmerie units is centralized, with the territory divided into gendarmerie companies.

The statistics are derived from 372 *Departmental Gendarmerie Companies* and 828 *Public Security Districts*, aggregated at the departmental level. This choice is based on two main reasons. First, security perimeters in France have been modified since

2011, with additions, mergers, and divisions, among other changes. Second, security policy in France is defined at the central level, then implemented at the departmental level before being applied to the 1,200 local and regional security perimeters. Although the division of French territory into 101 departments has remained unchanged since 2011, these statistics also cover overseas territories, such as *Saint-Martin*, *French Polynesia*, and *New Caledonia*, for the period 2012-2021. However, data for *Wallis and Futuna* is only available until 2016. For simplicity, in the remainder of this paper, these territories will be referred to as departments. Thus, statistics will be aggregated over 105 departments for the period 2012-2016 and 104 departments for the period 2017-2021. Finally, these data will be categorized based on the occurrence of a terrorist attack, a foiled attack, or the absence of such events in a department. An exception is noted in one department, where a terrorist attack occurred, followed by the foiling of another attack in 2016.

During the period 2012-2021, on average, 3.44% of french departments and territorial collectivities were affected by terrorist attacks. According to Figure 1, six departments were affected by these events in 2015, five departments in 2020, and four departments in 2016, 2017, and 2021. Three departments were impacted in other years, except for 2013, when only one department was concerned. Foiled attacks were recorded in only two departments in 2015, 2016, and 2021. Figure 1 also shows, for each year from 2012 to 2021, the number of departments that were affected by terrorist attacks (in gray), those that experienced foiled attacks (in orange), and those that both suffered an attack and successfully foiled others (in yellow).

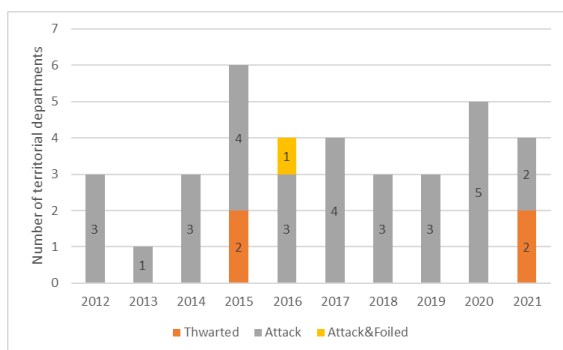


Figure 1: Evolution of events related to terrorist attacks in France (2012-2021).

In conclusion, although the average proportion of French departments affected by terrorism between 2012 and 2021 remained relatively low, certain years

such as 2015 and 2020 saw a higher concentration of attacks. This distribution indicates a persistent but geographically limited terrorist threat. Furthermore, the low number of recorded foiled attacks suggests either their rarity or a possible under-reporting or centralization of counter-terrorism efforts.

For the purpose of this study, we grouped the statistics into periods<sup>2</sup>, all starting from 2012. Figure 2 presents the number of observations per period as well as the percentage of events related to terrorist attacks.

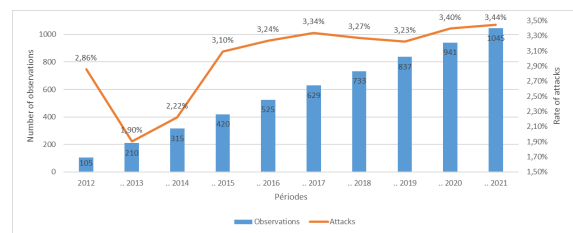


Figure 2: Evolution of terrorist attacks by periods starting from 2012.

Each period will be used as a training data. Our learning models will aim to predict terrorist-related events with the highest possible performance. It is important to note that these data are highly imbalanced. For example, between 2012 and 2013, 210 observations were recorded, of which 4 were related to terrorist attacks, accounting for 1.9%.

## 4 INTERPRETABLE LEARNING OF TERRORIST ATTACKS IN FRANCE

The objective of this work is to understand the evolution of terrorist behavior based on labeled data (by year and period). Initially, we propose to learn from the data corresponding to each year to deduce a distinct behavior for each year. Using a supervised learning method, we generate 10 distinct models, each corresponding to a year from 2012 to 2021. To ensure explainability and interpretability of the results, we will adopt a symbolic learning approach, based on techniques such as decision trees (Quinlan, 1993), decision rule generators (Ghosh et al., 2022), or Formal Concept Analysis (Meddouri and Maddouri, 2020). We propose to use well-known decision rules generators from the literature, implemented in *WEKA*<sup>3</sup>. Among the classifiers handling numerical and multi-

<sup>2</sup>Period 1: 2012. Period 2: 2012-2013. Period 3: 2012-2014 ... Period 10: 2012-2021

<sup>3</sup><https://ml.cms.waikato.ac.nz/weka>

Table 1: Number of Decision Rules Computed Annually by Generators.

| Year            | 2012 | 2013 | 2014 | 2015 | 2016 | 2017 | 2018 | 2019 | 2020 | 2021 | Avg.( $\pm$ Std.Dev.) |
|-----------------|------|------|------|------|------|------|------|------|------|------|-----------------------|
| ConjunctiveRule | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1 ( $\pm$ 0)          |
| DecisionTable   | 1    | 1    | 1    | 5    | 1    | 2    | 1    | 1    | 4    | 1    | 1,8 ( $\pm$ 1.12)     |
| DTNB            | 1    | 1    | 1    | 9    | 1    | 1    | 1    | 1    | 1    | 4    | 2,1 ( $\pm$ 1.76)     |
| FURIA           | 2    | 1    | 2    | 4    | 3    | 2    | 2    | 2    | 2    | 4    | 2,4 ( $\pm$ 0.76)     |
| JRIP            | 1    | 1    | 1    | 3    | 2    | 2    | 1    | 1    | 2    | 1    | 1,5 ( $\pm$ 0.6)      |
| MODLEM          | 8    | 3    | 7    | 6    | 6    | 2    | 6    | 7    | 8    | 6    | 5,9 ( $\pm$ 1.36)     |
| MOEFC           | 4    | 4    | 4    | 5    | 4    | 5    | 4    | 4    | 4    | 4    | 4,2 ( $\pm$ 0.32)     |
| NNge            | 6    | 3    | 6    | 5    | 5    | 2    | 5    | 4    | 6    | 6    | 4,8 ( $\pm$ 1.08)     |
| OLM             | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1 ( $\pm$ 0)          |
| OneR            | 1    | 1    | 1    | 2    | 2    | 2    | 2    | 2    | 2    | 2    | 1,7 ( $\pm$ 0.42)     |
| PART            | 4    | 1    | 1    | 4    | 2    | 2    | 2    | 2    | 3    | 4    | 2,5 ( $\pm$ 1)        |
| Ridor           | 1    | 1    | 1    | 3    | 1    | 2    | 1    | 1    | 1    | 1    | 1,3 ( $\pm$ 0.48)     |
| RoughSet        | 7    | 3    | 6    | 7    | 6    | 2    | 4    | 4    | 9    | 8    | 5,6 ( $\pm$ 1.88)     |
| ZeroR           | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1 ( $\pm$ 0)          |

class data, we mention *ConjunctiveRule* (Kalmegh, 2018), *DecisionTable* (Kohavi, 1995), *DTNB* (Hall and Frank, 2008), *FURIA* (Hühn and Hüllermeier, 2009), *JRIP* (Cohen, 1995), *Multi-Objective Evolutionary Algorithms for Fuzzy Classification*<sup>4</sup> (Jimenez et al., 2014), *NNge* (Martin, 1995), *OLM* (Ben David, 1992), *OneR* (Holte, 1993), *PART* (Frank and Witten, 1998), *Ridor* (Gaines and Compton, 1995), *RoughSet* (Wojna et al., 2023), and *ZeroR* (Sangeorzan, 2020). The choice of these generators is motivated by their interpretability and their relevance in the state-of-the-art of symbolic learning.

The selection of decision rule generators was based on two main criteria: interpretability and representation of state-of-the-art symbolic learning methods. Rule-based models are inherently interpretable because they express knowledge as human-readable IF-THEN rules, which is essential in security-sensitive domains such as counter-terrorism. To ensure diversity and relevance, we considered a set of generators implemented in WEKA, a widely recognized platform for benchmarking machine learning models. This selection includes classical symbolic learners such as *OneR*, *PART*, and *JRIP*, which serve as standard baselines for rule induction; advanced fuzzy and evolutionary approaches like *FURIA* and *MOEFC*, designed to handle uncertainty and multi-objective optimization; instance-based and hybrid methods such as *NNge* and *DTNB*, which combine rule induction with probabilistic reasoning; and rough set or formal concept-based methods like *RoughSet* and *MODLEM*, which are well-established in interpretable knowledge discovery. Together, these methods cover different paradigms deterministic, probabilistic, fuzzy, and evolutionary, while maintaining the interpretability requirement. Moreover, their ex-

tensive citation in recent literature on interpretable machine learning and decision support systems confirms their relevance as state-of-the-art techniques.

We acknowledge that standard baselines such as decision trees, logistic regression, and random forests are commonly used in predictive modeling. However, these methods were deliberately excluded from this study because our primary objective is to ensure interpretability and explainability, which are critical in security-sensitive contexts. While tree-based and ensemble methods often achieve higher predictive accuracy, they typically may lack the transparency required for operational decision-making. Future work will include these baselines to provide a broader comparison in terms of predictive performance versus interpretability.

In Table 1, we present the number of decision rules calculated by the previously mentioned generators for each year in the period 2012-2021. Regardless of the analyzed year, *ConjunctiveRule*, *OLM*, and *ZeroR* produce only a single decision rule to describe the annual behavior. *DecisionTable*, *DTNB*, *FURIA*, *JRIP*, *OneR*, *PART*, and *Ridor* generate very few decision rules (less than 3 on average). In contrast, *MOEFC*, *NNge*, *RoughSet*, and *MODLEM* generate an average of 4.2 ( $\pm$ 0.32), 4.8 ( $\pm$ 1.08), 5.6 ( $\pm$ 1.88), and 5.9 ( $\pm$ 1.36) decision rules per year, respectively.

In Table 2, we present the number of decision rules generated by the previously mentioned generators for different periods. Unlike the previous observations, the *MODLEM* and *NNge* generators produce a significantly higher number of decision rules, with an average of 26.4 ( $\pm$ 10) and 48.3 ( $\pm$ 31.41) rules, respectively. Similarly, *DTNB* generates an average of around a hundred decision rules (128.8 ( $\pm$ 111.6)), while *RoughSet* generates several hundred decision rules, with an average of 3376.8 ( $\pm$ 3940.07). The number of decision rules generated by these generators adapts to the size of the training data.

<sup>4</sup>In the rest of this article, the classifier *Multi-Objective Evolutionary Algorithms for Fuzzy Classification* will be abbreviated as *MOEFC*.

Table 2: Number of Decision Rules generated per Period.

| Year            | 2012 | ..13 | ..14 | ..15 | ..16 | ..17 | ..18 | ..19 | ..20  | ..21  | Avg.(±Std.Dev.)   |
|-----------------|------|------|------|------|------|------|------|------|-------|-------|-------------------|
| ConjunctiveRule | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1     | 1     | 1 (±0)            |
| DecisionTable   | 1    | 1    | 1    | 2    | 9    | 1    | 1    | 3    | 32    | 6     | 5,7 (±5.43)       |
| DTNB            | 1    | 1    | 1    | 22   | 23   | 111  | 155  | 254  | 309   | 411   | 128,8 (±111.6)    |
| FURIA           | 2    | 4    | 5    | 4    | 8    | 13   | 10   | 9    | 11    | 14    | 8 (±3.09)         |
| JRIP            | 1    | 2    | 2    | 3    | 1    | 2    | 2    | 1    | 3     | 2     | 1,9 (±0.49)       |
| MODLEM          | 8    | 11   | 15   | 18   | 25   | 29   | 34   | 39   | 42    | 43    | 26,4 (±10)        |
| MOEFC           | 4    | 4    | 4    | 4    | 4    | 4    | 7    | 7    | 5     | 8     | 5,1 (±1.21)       |
| Nnge            | 6    | 7    | 13   | 23   | 28   | 40   | 55   | 70   | 104   | 137   | 48,3 (±31.41)     |
| OLM             | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1     | 1     | 1 (±0)            |
| OneR            | 1    | 1    | 1    | 2    | 2    | 2    | 2    | 2    | 2     | 3     | 1,8 (±0.43)       |
| PART            | 4    | 3    | 8    | 5    | 6    | 9    | 8    | 16   | 15    | 12    | 8,6 (±3.2)        |
| Ridor           | 1    | 1    | 1    | 1    | 2    | 1    | 2    | 1    | 2     | 7     | 1,9 (±0.98)       |
| RoughSet        | 7    | 8    | 75   | 189  | 363  | 823  | 1619 | 2260 | 17855 | 10569 | 3376,8 (±3940.07) |
| ZeroR           | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1    | 1     | 1     | 1 (±0)            |

As shown in Tables 1 and 2, decision rule generators described in the literature produce models of differing sizes, measured by the number of decision rules generated. The generators *ConjunctiveRule*, *OLM*, and *ZeroR* generate only a single rule at a time, with an average deviation equal to 0. In contrast, the other rule generators exhibit highly variable average deviations.

In conclusion, the use of decision rule generators allows us to interpret and explain the generated learning models. For example, in the appendix of this paper, we present the learning models produced by *MODLEM*. For each year or period, a distinct learning model is obtained in the form of decision rules set, ensuring the model's explainability. In the case of *MODLEM*, it describes criminal behavior in 2012 through 8 decision rules, whereas for the following year, it requires only 3 decision rules (see Table 2). The analysis of these models will help to better understand and interpret the evolution of crime and threats, such as unrest, riots, terrorist attacks, and other phenomena. This analysis will consist of comparing the generated sets of rules in pairs to measure their similarity. If two sets of rules are highly similar, this may indicate that criminal behavior has changed little. Else, if the sets of decision rules are slightly or not at all similar, this will suggest a significant evolution in criminal behavior, specifying the differences between the rules.

## 5 EXPERIMENTAL STUDY

The purpose of this section is to study the performance of decision rule generators for predicting events related to terrorist attacks, using the previously generated learning models. To evaluate these performances, we rely on standard classification criteria, such as Error Rates, Recall/Sensitivity, ROC-Area

(AUC-ROC<sup>5</sup>), PRC Area (AUPR<sup>6</sup>), and Kappa Statistic.

These indicators will allow us to analyze the effectiveness of each decision rule generator to correctly predict the departments and periods associated with terrorist attacks, as well as their ability to avoid false positives and false negatives. Each generator will be evaluated based on its performance across different periods in order to test the robustness of the models against data imbalances. The performances will be compared among the generators to identify those that offer the best trade-offs between model complexity (number of decision rules) and prediction accuracy. Although the dataset is highly imbalanced, with terrorist-related events representing less than 4% of the observations, we deliberately avoided applying oversampling or synthetic data generation techniques such as *SMOTE* to preserve the interpretability and fidelity of the models. Instead, we addressed class imbalance at the evaluation stage by adopting metrics that are robust to skewed distributions, including Recall, AUPR, and Kappa statistic. This methodological choice ensures that the models remain explainable while still providing meaningful performance indicators for rare but critical events. The learning data is detailed in Table 3 by period.

The data used for testing generalization is described by 104 attributes. The number of observations includes 105 for the data from the years 2012 to 2016, and 104 from 2017 to 2021.

The experimental protocol consists of 9 experiments, the details of each one are presented in table 4. Each experiment is designed to test the ability of the generated learning models to predict events in the following year, using training data from successive periods.

<sup>5</sup>Area Under Curve - Receiver Operating Characteristic

<sup>6</sup>Area Under Precision-Recall

Table 3: Characteristics of the Learning Data.

| Year            | 2012 | ..2013 | ..2014 | ..2015 | ..2016 | ..2017 | ..2018 | ..2019 | ..2020 | ..2021 |
|-----------------|------|--------|--------|--------|--------|--------|--------|--------|--------|--------|
| Characteristics | 104  | 104    | 104    | 104    | 104    | 104    | 104    | 104    | 104    | 104    |
| Observations    | 105  | 210    | 315    | 420    | 525    | 629    | 733    | 837    | 941    | 1045   |

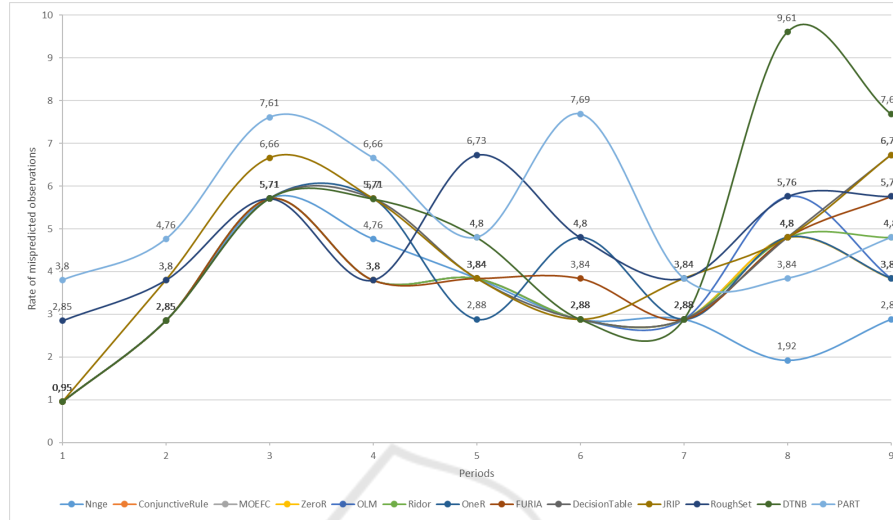


Figure 3: Evolution of mispredicted observations rates.

Table 4: Learning and generalization sets of data.

| Experimentation | Learning Data | Generalization Data |
|-----------------|---------------|---------------------|
| 1               | 2012          | 2013                |
| 2               | 2012 .. 2013  | 2014                |
| 3               | 2012 .. 2014  | 2015                |
| 4               | 2012 .. 2015  | 2016                |
| 5               | 2012 .. 2016  | 2017                |
| 6               | 2012 .. 2017  | 2018                |
| 7               | 2012 .. 2018  | 2019                |
| 8               | 2012 .. 2019  | 2020                |
| 9               | 2012 .. 2020  | 2021                |

More specifically, in each experiment, the training data covers a period from 2012 to a given year, and the generated learning models are then evaluated on generalization data corresponding to the following year. This approach allows testing the generalization of the learning models to events that occur after the training period, in order to assess their robustness and their ability to predict future trends in safety and crime.

### 5.1 Analysis of Mispredicted Observation Rates

According to Figure 3, the rates of incorrectly predicted observations by the decision rule generators remain relatively stable throughout the 9 experimental periods. This confirms that these generators, as supervised learning methods, offer stable prediction/classification performance. Among the decision

rule generators, *NNge* minimizes the misclassification rate, with an average of 3.19% ( $\pm 1.06$ ). It is closely followed by *ConjunctiveRule*, *MOEFC*, and *ZeroR*, which display similar performance, with an average of 3.51% ( $\pm 0.99$ ). The highest rates are observed for *PART* and *DTNB*, with 5.31% ( $\pm 1.34$ ) and 4.79% ( $\pm 2.13$ ), respectively. In summary, the average performance of the decision rule generators tested varies between 3.19% and 5.31%, highlighting a general stability in their prediction capabilities.

### 5.2 Recall/Sensitivity Rates Analysis

The Recall, or Sensitivity (Recall/Sensitivity) measure evaluates a learning model's ability to identify all actual positive observations. It indicates the proportion of true positives correctly classified as such. According to Figure 4, recall rates range between 0.9% and 0.99%, with an average of 0.96% and a very low average deviation ( $\pm 0.01$ ). This shows that the decision rule generators are able to identify most of the actual positive observations. This capability is even more important in the context of our application, where the cost of false negatives is high, both economically and sociologically.

### 5.3 Kappa Statistic Analysis

The Kappa Statistic measures the difference between the observed agreement and the agreement expected

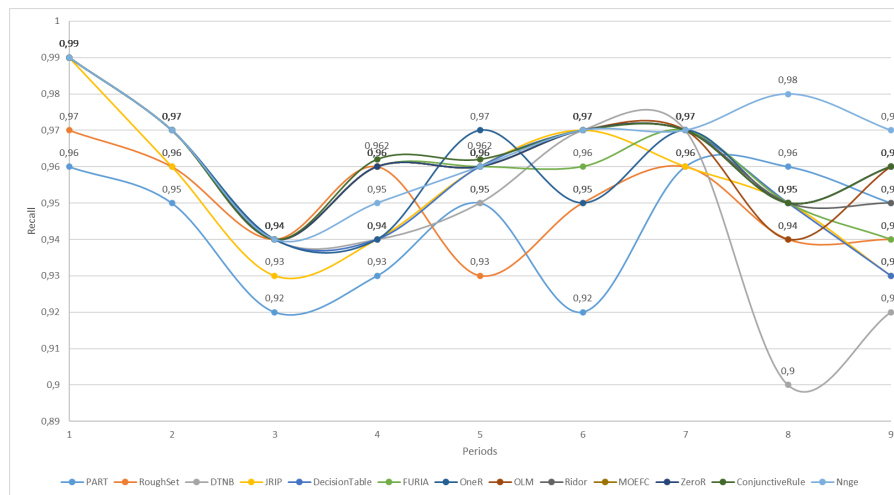


Figure 4: Evolution of Recall/Sensitivity Rates.

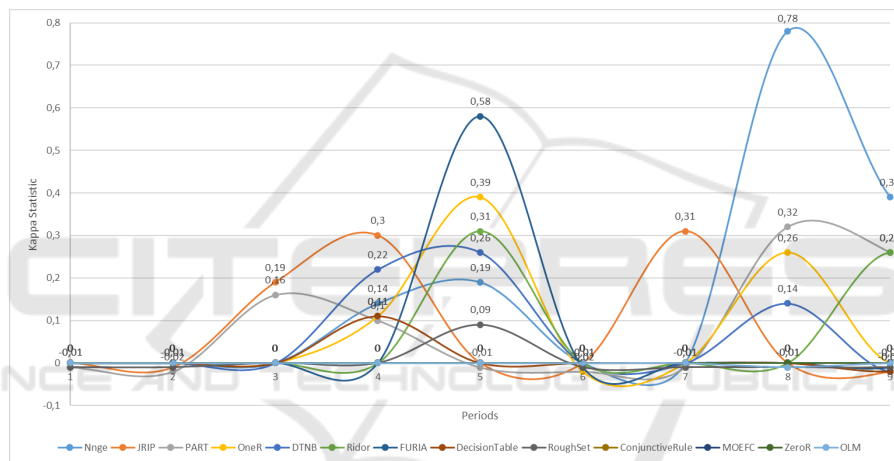


Figure 5: Kappa Statistic evolution.

by pure chance. When Kappa Statistic is close to 1, it indicates that the model performs much better than random chance. If Kappa Statistic is close to 0, it means that the model performs no better than a random prediction. A negative Kappa Statistic suggests that the model performs worse than random chance. According to Figure 5, most decision rule generators have a Kappa value around 0.05 ( $\pm 0.07$ ), except for *NNge*, which reaches an average of 0.17 ( $\pm 0.19$ ).

#### 5.4 ROC Area (AUC-ROC) Analysis

The ROC curve (Receiver Operating Characteristic) describes the evolution of Sensitivity (or true positive rate) as a function of 1 minus Specificity (anti-specificity) as the decision threshold changes. The term ROC comes from the intercommunication between systems, where these curves are used to analyze a model's ability to separate the signal from the

background noise. The area under the ROC curve, or AUC (Area Under the Curve), measures the area under the ROC curve, which plots the true positive rate against the false positive rate for different classification thresholds. It allows us to evaluate a model's ability to distinguish between positive and negative classes. The AUC is also useful for comparing model performances at different thresholds: a value of 1 indicates perfect classification, while a value of 0.5 suggests random performance. According to Figure 6, the decision rule generators *ConjunctiveRule*, *MOEFC*, *ZeroR*, and *OLM* show a constant evolution throughout the experimental periods, with a stable rate of 0.5 and zero average deviation. In contrast, the generators *DTNB*, *DecisionTable*, and *FURIA* achieve results above 0.5, with respective averages of 0.64 ( $\pm 0.15$ ), 0.62 ( $\pm 0.14$ ), and 0.61 ( $\pm 0.13$ ). In summary, most of the rule generators experimented with produce learning models whose performance is

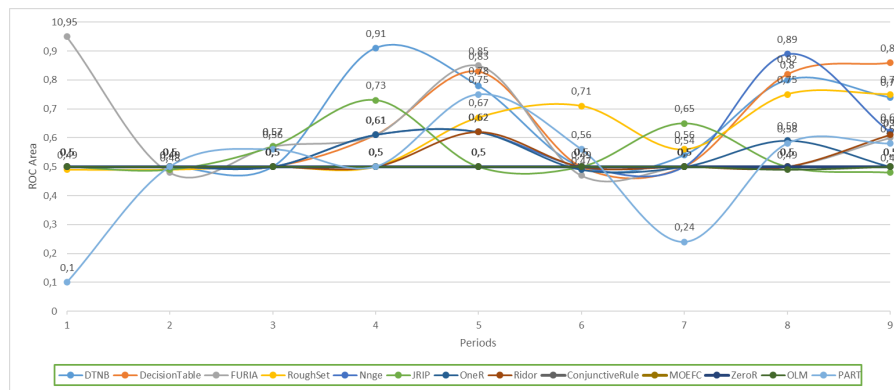


Figure 6: ROC Area (AUC-ROC) evolution.

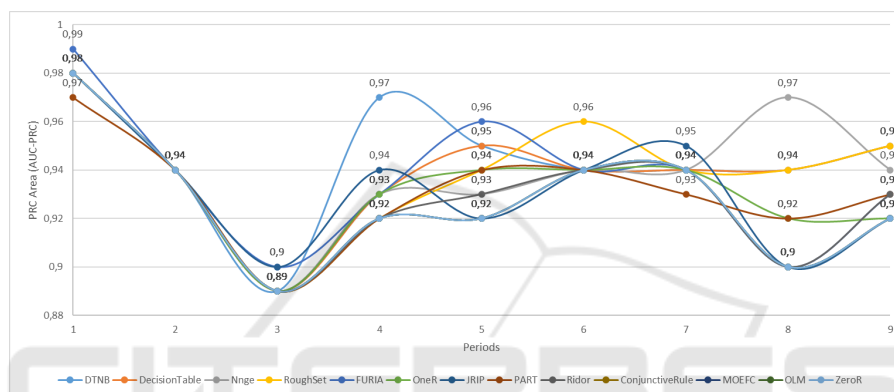


Figure 7: PRC Area (AUPR) evolution.

close to random classification, with an average of  $0.55 (\pm 0.07)$ , except for *DTNB*, *DecisionTable*, and *FURIA*. Although ROC-AUC values are close to 0.5, this is expected under severe class imbalance and does not fully reflect the models' ability to identify rare positive events. Therefore, we emphasize Recall and AUPR as more relevant metrics for this context.

### 5.5 PRC Area (AUPR) Analysis

The Precision-Recall Curve (PRC) often complements the ROC curve. It describes the evolution of precision as a function of Recall as the decision threshold changes. To summarize this curve, we use the area under it, called AUPR (Area Under the Precision-Recall Curve). The AUPR is especially useful when there is an imbalance between the classes, as is the case in our study. A higher score indicates better performance in identifying the positive class. According to Figure 7, the decision rule generators are near 1, with an average of  $0.93 (\pm 0.02)$ . This suggests that all rule generators perform well in terms of Precision and Recall. The generators *DTNB*, *NNge*, *RoughSet*, and *FURIA* achieve an average of  $0.94 (\pm 0.02)$ , placing them among the best models in

terms of precision-recall.

## 6 CONCLUSION AND PERSPECTIVES

This study evaluated the performance of decision rule generators in predicting terrorist-related events in France using data from 2012 to 2021. While most generators showed stable performance, some stood out in terms of precision and efficiency. Notably, *NNge*, *MOEFC*, and *RoughSet* achieved high precision-recall and AUPR scores (approaching 1), indicating strong capabilities despite class imbalance. In contrast, *DTNB* and *DecisionTable* yielded lower AUC and Kappa scores, reflecting weaker discrimination in imbalanced contexts. Low Kappa and RAE values for certain models suggest limited but acceptable agreement between predictions and actual outcomes.

Although ROC-AUC and Kappa values remain modest, limiting real-world predictive utility, this is largely due to extreme class imbalance and the choice to preserve interpretability by avoiding oversampling

and black-box models. Nevertheless, the models achieve high Recall ( $\approx 96\%$ ) and AUPR ( $\approx 0.93$ ), crucial for minimizing false negatives in security-sensitive applications. These findings highlight the value of interpretable decision-support tools, even with limited discriminative power.

To further investigate model interpretability, we propose analyzing rule set similarities between *MOEFC* and *NNge*, which generate an average of  $4.2(\pm 0.32)$  and  $4.8(\pm 1.08)$  rules respectively, balancing simplicity and performance. Conversely, methods like *OLM* and *Ridor*, with fewer rules, offer lower complexity. Interestingly, generators such as *ConjunctiveRule*, *OLM*, and *ZeroR* maintain good performance despite minimal spatial and temporal complexity. Exploring sequential or parallel rule generation could enhance robustness while managing complexity, offering a promising trade-off between explainability and performance.

Finally, while decision rule models are inherently interpretable, understanding the generated rules is essential to link predictions with underlying societal and political factors. Enhancing model interpretability can strengthen trust and support informed decision-making in real-world scenarios.

## REFERENCES

- Ben David, A. (1992). Automated generation of symbolic multiattribute ordinal knowledge-based dss: Methodology and applications. *Decision Sciences*, 23:1357–1372.
- Cohen, W. W. (1995). Fast effective rule induction. In *Twelfth International Conference on Machine Learning*, pages 115–123. Morgan Kaufmann.
- Da Silva, A. R. C., de Paula Júnior, I. C., da Silva, T. L. C., de Macêdo, J. A. F., and Silva, W. C. P. (2020). Prediction of crime location in a brazilian city using regression techniques. In *Proceeding of IEEE 32nd International Conference on Tools with Artificial Intelligence (ICTAI'20)*, pages 331–336.
- Frank, E. and Witten, I. H. (1998). Generating accurate rule sets without global optimization. In Shavlik, J., editor, *Fifteenth International Conference on Machine Learning*, pages 144–151. Morgan Kaufmann.
- Gaines, B. and Compton, P. (1995). Induction of ripple-down rules applied to modeling large databases. *Journal of Intelligent Information Systems*, 5:211–228.
- Ghosh, B., Malioutov, D., and Meel, K. S. (2022). Efficient learning of interpretable classification rules. *Journal of Artificial Intelligence Research*, 74:1823–1863.
- Hall, M. and Frank, E. (2008). Combining naive bayes and decision tables. In *Proceedings of the 21st Florida Artificial Intelligence Society Conference (FLAIRS)*, pages 318–319. AAAI press.
- Holte, R. (1993). Very simple classification rules perform well on most commonly used datasets. *Machine Learning*, 11:63–91.
- Hühn, J. and Hüllermeier, E. (2009). Furia: An algorithm for unordered fuzzy rule induction. *Data Min. Knowl. Discov.*, 19:293–319.
- Jimenez, F., Sánchez, G., and Juarez, J. (2014). Multi-objective evolutionary algorithms for fuzzy classification in survival prediction. *Artificial Intelligence in Medicine*, 60.
- Kalmegh, S. R. (2018). Comparative analysis of the weka classifiers rules conjunctive rule and decision table on indian news dataset by using different test modes. *International Journal of Advanced Research in Computer Science*, 7(2):01–09.
- Kohavi, R. (1995). The power of decision tables. In *8th European Conference on Machine Learning*, pages 174–189. Springer.
- Martin, B. (1995). Instance-based learning: Nearest neighbor with generalization. Master's thesis, University of Waikato, Hamilton, New Zealand.
- Meddouri, N. and Beserra, D. (2024). Apprentissage interpretable de la criminalité en france (2012-2021). In *Proceeding of the EGC workshop Gestion et Analyse de données Spatiales et Temporelles with the 24<sup>ième</sup> Journées Francophones en Extraction et Gestion des Connaissances (GAST/EGC 24)*, pages 41–43.
- Meddouri, N. and Maddouri, M. (2020). Efficient closure operators for fca-based classification. *International Journal of Artificial Intelligence and Machine Learning*, 10:79–98.
- Mucchielli, L. (2008). Une société plus violente ? Une analyse socio-historique des violences interpersonnelles en France, des années 1970 à nos jours. *Deviance et Société*, 32(2):115–147.
- Quinlan, J. R. (1993). *C4.5: programs for machine learning*. Morgan Kaufmann Publishers Inc., San Francisco, CA, USA.
- Saidi, F. and Trabelsi, Z. (2022). A hybrid deep learning-based framework for future terrorist activities modeling and prediction. *Egyptian Informatics Journal*, 23(3):437–446.
- Sangeorzan, L. (2020). Effectiveness analysis of zero and j48 classifiers using weka toolkit. *Bulletin of the Transilvania University of Brasov. Series III: Mathematics and Computer Science*, pages 481–486.
- Tolan, G., Abou-El-Enien, T., and Khorshid, M. (2015). Hybrid classification algorithms for terrorism prediction in middle east and north africa. *International Journal of Emerging Trends & Technology in Computer Science (IJETTCS)*, 4:23–29.
- Wojna, A., Latkowski, R., and Kowalski, L. (2023). *Rses-lib: User Guide*. Accessed: 2025-01-04.