

# Epileptic Seizure Detection Using EEG Signals and Extreme Gradient Boosting

Paul Vanabelle<sup>\*a</sup>, Pierre De Handschutter<sup>b</sup>, Riëm El Tahry<sup>c</sup>, Mohammed Benjelloun<sup>b</sup>, Mohamed Boukhebouze<sup>a</sup>

**Abstract**—The problem of automated seizure detection is treated using clinical electroencephalograms (EEG) and machine learning algorithms on the Temple University Hospital EEG Seizure Corpus (TUSZ). Performances on this complex dataset are still not encountering expectations. The purpose of this work is to determine to which extend the use of larger amount of data can help to improve the performances. Two methods are explored; a standard partitioning on a recent and larger version of the TUSZ; and a Leave-One-Out approach used to increase the amount of data for the training set. XGBoost, a fast implementation of the gradient boosting classifier, is the ideal algorithm for these tasks. The performances obtained are in the range of what is reported until now in the literature with deep learning models. We give interpretation to our results by identifying the most relevant features and analyzing performances by seizure types. We show that generalized seizures tend to be far more well predicted than focal ones. We also notice that some EEG channels and features are more important than others to distinguish seizure from background.

## I. INTRODUCTION

One of the most common way to diagnose epileptic seizure is to measure the electrical activity of the cerebral cortex by performing a non-invasive electroencephalogram (EEG) [1]. It is a relatively unexpensive and easy way to proceed compared to other techniques such as MRI or intrusive methods [2]. Consequently, most of works to automatically analyze seizures have been done on these signals through time series processing methods.

EEG-based seizure detection has been extensively studied, however it has rarely reached performance that could durably help neurologists. This is due to several factors. The complexity of the epilepsy and of the brain waves is the first one. The EEG signals are non-stationary and the statistical features of these signals are different between patients and over time. A second factor is the quality of the datasets used for the machine learning task. An ideal dataset needs to be a good representation of all the varieties of EEG signals that could appear during epileptic seizure. Even if high performances were obtained with some of the publicly available datasets, they have typically inherent restrictions.

For instance, they can be too specific, regarding a certain type of population as CHB-MIT dataset [3]. Indeed, this dataset is only focused on teens and children and, as it is generally recognize that EEG signals differ with age [4], it is potentially difficult to generalize the resulting work on adults. It is also particularly more dedicated to the prediction tasks thanks to the long duration of the recordings. Good results are presented in [5] on this task. Datasets can also be too small or incompletely documented such as Bonn dataset where there is no indication of the type of seizure and patients features [6]. With this one, very high performance with various deep learning architectures using raw data were achieved, but that seems hard to generalize on more sophisticated datasets [7], [8].

Therefore these experiments do not correspond to clinical conditions and the results are not representative of current clinical performances. This state of fact initiated the development of the Temple University Hospital EEG Seizure Corpus (TUSZ). Up to now it is the largest open source corpus of this kind, and represents an accurate characterization of clinical conditions [9]. TUSZ Corpus is a complex dataset offering a great diversity in terms of patients, seizures or EEG montages (the placement of the electrodes). Machine learning models have been proposed in [9] and [10]. These works are based on deep architectures but the results are not yet convincing. Some architectures obtained good specificities but sensitivities are still low resulting in a low rate of false alarms but in a lot of seizure episodes not detected. The first versions of the dataset were used for these works. They were already offering more patients for the studies than the other seizures corpora, the latest versions providing even more.

All these observations led us to the following question: **"Does the increase of data lead to better results?"**. We saw that a common difficulty for a generic modeling is the lack of data, both in quantity and quality, for having relevant results through machine learning and for having models that generalize well on the data coming from new patients.

To increase the amount of training data, we first work on a more recent release of the TUSZ offering a larger training set, with a classic approach to machine learning and data partitioning. This give the first method. The second method is based on a Leave-One-Out approach with iterations, in the manner of a cross-validation, to extend the training set with the data of the test set. It means that, at each iteration, we keep a patient of the initial test set for the testing phase and perform the training and validation of the model on

<sup>\*</sup> Corresponding author: paul.vanabelle@cetic.be

<sup>a</sup> Data Science Department, Centre of Excellence in Information and Communication Technologies (CETIC), Avenue Jean Mermoz 28, 6041 Charleroi, Belgium

<sup>b</sup> Computer Science Unit, Faculty of Engineering, University of Mons, rue de Houdain 9, 7000 Mons, Belgium

<sup>c</sup> Refractory Epilepsy Centre, University Hospital of Saint-Luc, Avenue Hippocrate 10, 1200 Brussels, Belgium

the rest of the data. XGBoost is the selected algorithm for our classifier. It is a fast implementation of the gradient boosting classifier which allows us to perform quickly the multiple iterations required by the second method. As it is a decision tree based algorithm, it has also the advantage to be interpretable; the measure of the features importance on the training set is returned by the algorithm. As inputs of our algorithm, temporal and frequential features are used with some others processed by a specific python library for EEG signals, called pyEEG. The methods are applied to the 2 most important subsets of the TUSZ in term of montage leading to 4 experiments.

As a result, with a fast computing approach, we managed to present results that matched the best in the Temple University Hospital (TUH) literature [9]. We analyze the performances by type of seizure and we show that the seizures which are categorized as generalized are easier to recognize and in which proportion. Finally the interpretable property of the XGBoost algorithm shows the features and the EEG channels that are the most discriminant to distinguish a seizure from the background, or ictal from interictal activity. Interpretability is important as emphasized in [11], as much as the collaboration between data scientists and clinicians in the early steps of such a study. We therefore report in the discussion section some ideas for reflection from our meetings with neurologists.

The remaining of the paper is organized as follows. In the next Section, we present the state-of-the-art on the TUSZ corpus and the use of gradient boosting algorithm in the context of seizure detection. In Section III, we describe the TUSZ dataset. We also introduce the two methods used for our experiments. We will close the Section III by the definition of the metrics considered in this work. Experimental analysis, results and discussion are presented in Section IV. We close this paper with the conclusion and the futures perspectives in Section V.

## II. RELATED WORK

As the first official release of the TUH EEG Seizure Corpus was done in April 2017, there is not a lot of available literature on machine learning techniques applied on this dataset. Most of the experiments come from Prof. Picone's team at Temple University and have led to Deep Learning approaches.

Several architectures are extensively described in [9] and the results presented were the first reported on this dataset. In this paper, the authors report a number of 64 patients on whom algorithms were trained and a number of 50 patients for the testing. Various architectures are applied, including Hidden Markov Models (HMM), Long Short Term Memory (LSTM), Convolutional Neural Networks (CNN) combined with MultiLayer Perceptron (MLP) and an hybrid CNN/LSTM model. Dimensionality reduction is usually performed, through (Incremental) Principal Component Analysis ((I)PCA). The results are based on metrics that are calculated permissively with the Any Overlap method (OVL); it consists in considering a totally good detection if

a detected event touches at least a period of time (a segment) of the real event, instead of calculating the metrics segment by segment. Specificity varies between 70 and 80% except in the hybrid CNN/LSTM model where an outstanding specificity of 96.86% is reached, resulting in a very low false alarm rate of 7 FA/24h. On the other hand, sensitivity is always low, 30% for this last model and up to 40 % for the others with a false alarm rate at minimum 77 FA/24h.

A similar work uses Gated Recurrent Units (GRU), a particular kind of Recurrent Neural Networks (RNN) that is faster but less accurate than LSTM [10]. A CNN/GRU approach is compared to the CNN/LSTM approach and shows the same sensitivity (30%) and a slightly lower specificity (91%). Initialization techniques and regularization methods are also discussed. Although version 1.1.1 of TUSZ was reported to be used in this case, with 196 patients for training and 50 for testing, the results remained in the same order of magnitude compared to the first paper. Hence, none of as of today state of art models have been able to reach satisfying sensitivities and, except a deep complex CNN/LSTM architecture with many layers, specificities and false alarm rates are not outstanding neither.

Gradient Boosting is an ensemble method using boosting principles. In short, it fits a classifier (generally a decision tree) to the data given in input and calculates residuals. A new classifier is then adjusted on these residuals. The procedure is repeated while the validation error decreases. A great advantage of using an ensemble method based on decision trees is that estimates of feature importance can be provided from a trained predictive model. This property can give significant insights about the features and on the analysis process.

XGBoost library is an efficient and distributed implementation of the Gradient Boosting algorithm [13]. The main strength of XGBoost is its scalability which allows parallel and distributed computing and makes learning and model exploration faster. Moreover, several other improvements were added such as techniques for reducing overfitting giving to XGBoost better performances than the generic boosting algorithm. We used the XGBoost library through the python package for this work.

Although Gradient Boosting is known to reach high performances in several tasks and is easier to parameterize than deep learning architectures, there was any use of it on "Big Data" seizure detection corpus such as TUSZ, to the best of our knowledge. However, we have to mention [14] where a Gradient Boosting approach is used on Freiburg dataset (which is not publicly available anymore) made of 21 patients. Their approach is patient-specific in the sense that they built a custom model for each patient. The data of one patient is actually split in a train subset and a test subset, before feeding a model. The procedure is repeated for all the other patients, resulting in as many models as patients. In the present work, we use the opposite approach which consists of a generic modeling.

Concerning the interpretable property we are trying to exploit by means of xgboost, we must mention that several

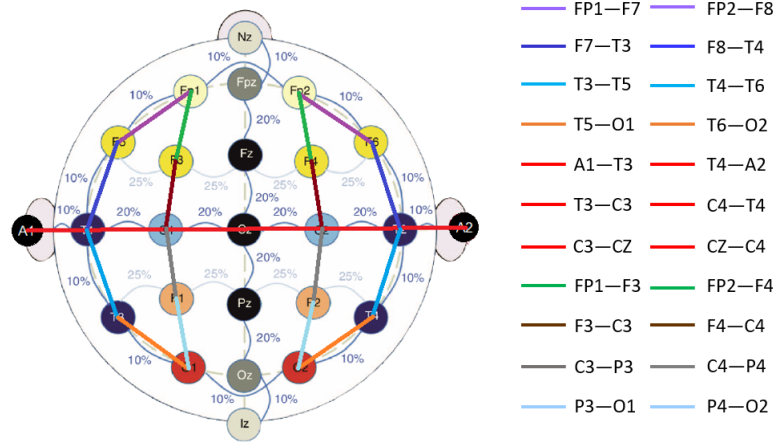


Fig. 1: Placement and connections of electrodes in a 10/20 placement with bipolar TCP montage [12]

other works approach the problem in a similar way [15], [16]. Both works consider a generic modeling on multiple patients and use the Bonn dataset for this purpose. The first one uses decision tree-based modelization, including a random forest classifier. The conclusion emphases the good performance of the random forest algorithm which is an ensemble method similar to gradient boosting. The second tries to establish rules from a fuzzy logic system. The interpretability of the resulting fuzzy rules is discussed. Bonn dataset only provides one EEG channel which limits the scope of interpretation.

### III. MATERIAL AND METHODS

#### A. Dataset

Temple University Hospital EEG Seizure Corpus (TUSZ) is a subset of the larger EEG signals database from the Temple University Hospital (TUH). TUSZ dataset is remarkable by the high number of patients and configurations it contains. Indeed, in the considered version 1.2.1 of this subset released in April 2018, there are 266 patients in the training set (of whom 118 suffering from seizures), 50 in the test set (of whom 38 suffering from seizures) and more than 40 different configurations of electrodes.

TUSZ dataset was established by identifying the sessions in the TUH EEG Corpus that were the most likely to contain seizure events. Three annotation tools were used to identify events of clinical interest: Natural Language Processing (NLP), a commercial software, and a three-pass system [17]. Then an annotation team manually annotated the data and showed that the NLP approach had been the most efficient as first step, revealing again the difficulty to automatically detect seizure events in a large dataset of EEG [18].

This dataset contains a huge diversity of configurations, in terms of sampling rate, references or montages (placement and connections between the electrodes). However two major montages predominate in the dataset : "Average Reference" (AR) for which the potentials are measured relatively to the average value of a subset of electrodes and "Linked Ears" (LE) for which the reference is located on the ear, electrically quiet. More statistics about TUH dataset are available in [18]

and [19], and more considerations on the montages and their consequences in [20].

Nevertheless, all recordings contain the electrodes of the standard 10/20 placement. We call "channels" the derivations of electrodes obtained with the TCP bipolar montage also called "double-banana". There are shown on Fig. 1, which comes from [12].

The recordings are split in six file folders : three theoretically for training purpose and the three remaining for the test. Each folder contains one specific montage, either AR, LE or a slight variant of AR. This variant is less represented and was not considered for this study.

For each patient, there are one or several sessions themselves containing files related to one or more recordings. There are five files for each recording :

- the *.edf* file contains the raw EEG signals and a header containing useful information (frequency, duration, date,...)
- the *.lbl* file contains event-based annotations that catch the beginning and end of each event on every channel of the 22 shown on figure 1. An event is either background or a particular type of seizure. Practically those annotations may indicate if an ictal phase begins or ends earlier or later on some channels
- the *.lbl.bi* is very similar to the previous one except that the events are just binary (*bckg* and *seiz*)
- the *.tse* file contains label-based annotations that catch the beginning and end of each event globally for all the channels. These are the most frequently used in Machine Learning research as they allow to give a single overall target for each time step
- the *.tse.bi* file is very similar to the previous one except that the events are just binary (*bckg* and *seiz*)

For each session, there is also a *.txt* file that contains general information about the patient, the recordings, in a non-structured way.

Epileptic seizures are usually classified in two main categories:

- The generalized crisis which begins bilaterally and

| Seizure Code | Seizure Type             | AR_Train | AR_Test | LE_Train | LE_Test |
|--------------|--------------------------|----------|---------|----------|---------|
| FNSZ         | Focal Non-Specific       | 13088    | 23301   | 10186    | 3418    |
| GNSZ         | Generalized Non-Specific | 6743     | 14226   | 10230    | 122     |
| SPSZ         | Simple Partial           | 1327     | 364     | 167      | 0       |
| CPSZ         | Complex Partial          | 8551     | 8114    | 5811     | 784     |
| ABSZ         | Absence                  | 14       | 0       | 479      | 339     |
| TNSZ         | Tonic                    | 424      | 857     | 0        | 0       |
| CNSZ         | Clonic                   | 0        | 0       | 0        | 0       |
| TCSZ         | Tonic Clonic             | 795      | 1335    | 938      | 2343    |
| ATSZ         | Atonic                   | 0        | 0       | 0        | 0       |
| MYSZ         | Myoclonic                | 0        | 1178    | 0        | 0       |

TABLE I: Seizure segments distribution by type in the different folders

synchronously from the outset in both cerebral hemispheres.

- The partial or focal crisis which starts in a localized part of the brain. It can be accompanied by a loss of consciousness; the abnormal electrical activity that started from focused location in the brain can gradually extend to the entire brain; and the crisis may end in convulsions.

However TUSZ is following a more refined classification, as it can be found in [21]. On this dataset 10 different types of seizures are defined, though some of them are less represented. Several types of seizures are actually covered by the generalized class or by the focal one. Among the generalized seizures, we find the "absence" which results in a momentary alteration of consciousness, the "tonic crisis" characterized by a strong variation of muscle tone, the "clonic crisis" characterized by jerks and the generalized "tonic-clonic" crisis characterized by convulsions, and also the atonic and myoclonic seizures. Among the focal seizures, we can specifically distinguish the complex and simple partial seizures if respectively there's a loss of consciousness or not. In table I, we have extracted from the dataset all the information about the distribution in terms of segments of each types of seizures in the several folders. The segments have in the present case a duration of one second. We can note in this table that non-specific seizures are dominant.

### B. Motivation of the proposed approach

Since we work on the TUSZ dataset, we tackled the problem of seizure detection through different angles. A prior work was based on the transformation of EEG signals into features and on the comparison of several classification models. We noticed during that experiment the good performances of ensemble learning methods as random forest and gradient boosting. Moreover the use of high-performance version of gradient boosting like XGBoost allowed to considerably reduce the training time, and thus to be more focused on fast iterations and optimization of parameters and hyper-parameters. We also noticed that the results were good when training and test sets were built after mixing non-overlapping ictal and inter-ictal segments on the whole population of recordings. Unfortunately, this method led to a kind of data leakage as a segment of one recording is not really different

from its direct neighbour of the same state. Hence, results were a bit less impressive when we tried to classify EEG segments on new patients. This work was using the first version of the dataset and limited to 59 patients for the training.

Interested by the abilities and the promises of some Deep Learning algorithms to directly extract features from complex time-series without preprocessing, we investigated that field without the expected success, as dealing with raw data is really resources-demanding. Another attempt by using EEG features with a Deep Learning algorithm (LSTM) came crashing down on the wall of bad performances related by the TUH literature [9].

Based on these previous approaches, we made several observations:

- The EEG features approach is still relevant for performance issues
- XGBoost is efficient and allows to quickly iterate
- There is a bad generalization on new patients in the test set

Taking up the last point, there are two ways to solve this issue. Either we move backward and consider a patient-specific modeling, or we increase the amount of training data for a generic modeling. The second choice is more relevant and it would be indeed interesting to see to which extent the increase of the amount of data would help to improve performances. EEG features and XGBoost will be the baseline of the two methods introduced hereafter. The first method aims to evaluate the performance of a classic split in terms of train, validation and test sets on the version 1.2.1 of TUSZ. With the second method, we search to increase the training dataset by replacing the classic split into a Leave One Out approach. It is a common technique used to increase a training set when there is a lack of data.

These two methods will be firstly applied to the data of AR Montages, in order to get significant results in reasonable time. Then experiments will be extended to both AR and LE.

### C. First Method : EEG Features + standard partitioning + XGBoost

The method is presented in Fig. 2. We first apply a preprocessing step to extract features from the raw EEG signals. These features are temporal and frequential. We

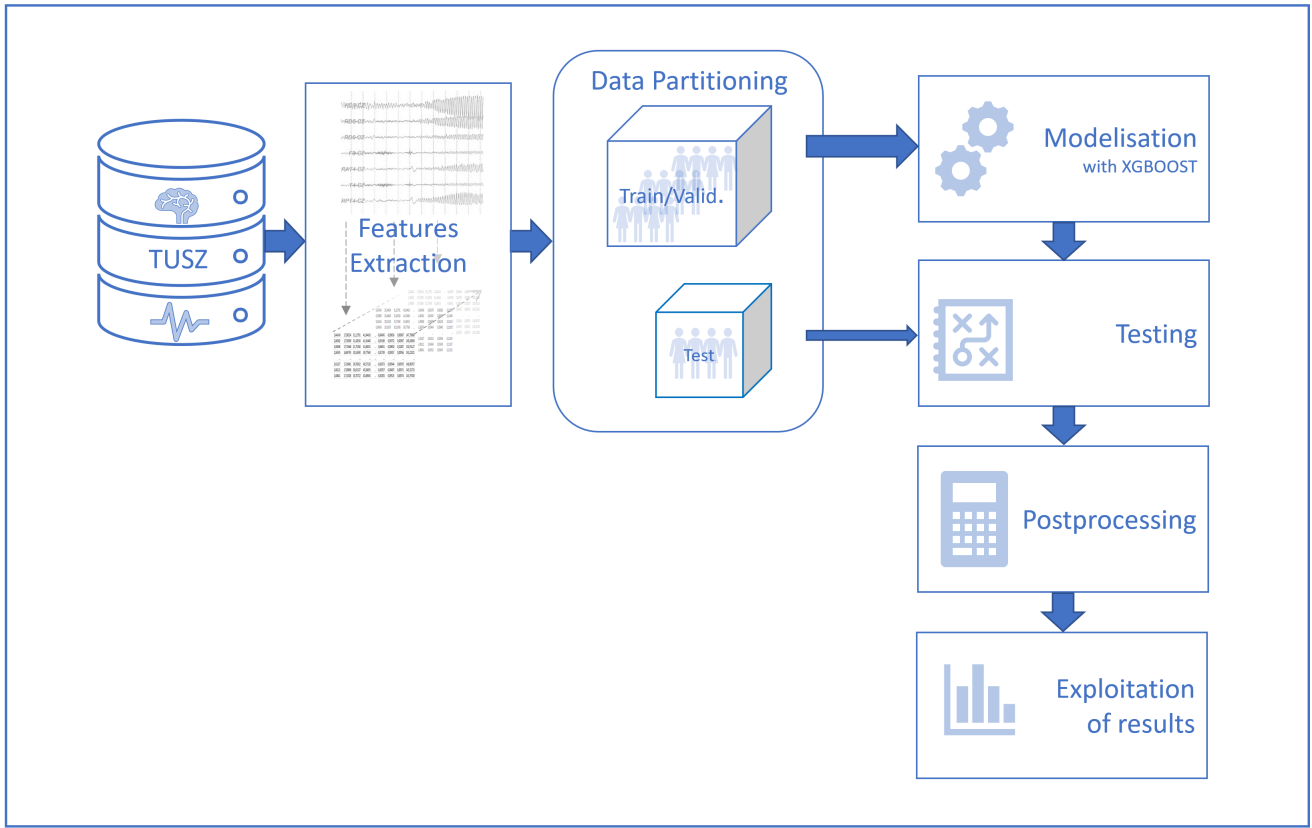


Fig. 2: First Method based on a classical train/test split

based ourselves on the literature review and Python's library PyEEG [22]. The following 22 features are then considered:

- Statistical parameters : mean, variance, skewness, kurtosis, interquartile range, minimum and maximum
- Other temporal features : Hjorth complexity and mobility, Petrosian fractal
- Frequential features : spectral density in the five frequency rhythms (alpha, beta, gamma, delta, theta) and the corresponding ratios ; spectral centroid and monotony. The five frequency ranges are defined in [23]. Delta range is the lowest : between 1 and 4 Hz, theta is between 4 and 8 Hz, alpha between 8 and 13 Hz, beta between 13 and 35 Hz and gamma > 35 Hz. Note that gamma rhythm is rarely modified by a seizure.

We split the raw EEG signal in non-overlapping 1s segments and compute these features for each of the 22 channels in the bipolar montage, resulting in a  $22 \times 22 = 484$  features vector for each segment. This preprocessing step is computed offline and each features matrix  $n\_segments \cdot n\_features$  is stored by recording.

After this preprocessing step, the segments of the TUSZ training set are randomly split in 80% of train and 20 % of validation. Then XGBoost with 400 estimators of depth 3 is trained. The ratio parameter is used to take into account class imbalance between non-seizure and seizure segments. To avoid problems of speed and memory, we have chosen to work with "gpu\_hist" as tree method, it's the XGBoost fast

histogram algorithm for GPU. The internal evaluation metric for the validation data was set to "AUC".

Once the model converged, the model is tested on all the recordings of the test set. For each recording, the model tries to predict the class of each segment. A smoothing function is then applied on those predictions to avoid isolated and improbable seizure/non-seizure states. This function is a sliding window on 19 segments applying a simple majority vote. At the end of the test set classification, we have confusion matrices for all the recordings and we can compute confusion matrices by patient or by types of seizures. Similarly, an overall confusion matrix is computed by summing all the recordings confusion matrices.

---

**Algorithm 1** Leave One Out Cross Test With XGBoost

---

**Input:** train\_patients\_indexes, test\_patients\_indexes, X, y

```

1: for i in test_patients_indexes do
2:   LeaveOneOutPatient= (X, y)[i]
3:   other_patients= (X,y) \ LeaveOneOutPatient
4:   XGBoost.fit (other_patients, validation_split=0.2)
5:   for record in LeaveOneOutPatient do
6:     pred=XGBoost.predict (X[record])
7:     pred=smooth(pred)
8:     Confusion_matrix (y[record], pred)
9:   end for
10: end for
  
```

---

#### D. Second Method : EEG Features + LOOCT + XGBoost

The second method is similar to the first one except that we apply a specific partitioning technique to increase the training data. This method, detailed in Algorithm 1, is close to a cross-validation method called Leave One Out Cross Validation (LOOCV). Actually, one should rather call it in this case LOOC-Test as explained below, as a single patient is left out for the test at each iteration. Similarly to a cross-validation method, we combine at the end the test results from the multiple rounds to come up with an estimate of the model's predictive performance. This second method is summarized on Fig. 3.

For each recording, we have a matrix of size  $n\_segments \cdot n\_features$  and the binary labels for each time segments (seizure or background). In the algorithm, one can see  $X$  as a list of length  $number\_of\_patients$  where each component is made of these features matrices for all the recordings of that patient, while  $y$  is the corresponding binary target vector. We also need to store separately the IDs of the patients in the TUSZ train sets and test sets.

Then, the algorithm can be applied. At each iteration, one patient is left out of the training data and is kept for the test. The segments of the other recordings are randomly split in 80% of train and 20 % of validation. Then similarly to the first method, a XGBoost model is trained, conserving the same parameters.

Once the model converged, the model is tested on each recording of the patient left away. Thus, for each recording, the model tries to predict the class of each segment. The smoothing function is still applied as post-processing.

The number of iterations of the LOOCT method is equal to the number of patients in the original test set present in the TUSZ. These are the only patients that are "left out" by the algorithm in order to have a kind of comparison to the first method in term of population. At the end of all the iterations, we have confusion matrices for all the recordings and we can compute a confusion matrix by patient or by type of seizure. Similarly, an overall confusion matrix is computed by summing all the recordings confusion matrices.

#### E. Metrics

The seizure detection is a classification problem, thus the evaluation of the models is based on confusion matrices (see Fig. 5). The number of true positives (TP) is in this case the number of EEG segments correctly classified as seizure; the true negatives (TN) are the number of EEG segments correctly classified as background; the false positives (FP) are the number of EEG segments classified as seizure that are actually background (False Positives are also known as false alarms in the medical field); the false negatives (FN) are the number of EEG segments classified as background though it is actually seizure.

|              |            | Predicted class      |                      |
|--------------|------------|----------------------|----------------------|
|              |            | Background           | Seizure              |
| Actual class | Background | True Negatives (TN)  | False Positives (FP) |
|              | Seizure    | False Negatives (FN) | True positives (TP)  |

Fig. 5: Confusion matrix for seizure detection purpose

The confusion matrices are used to compute some interesting ratios. Specificity, also known as true negative rate, is defined by

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

Sensitivity (or Recall) is the probability to detect an element of the positive class, it's defined by

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

Precision is the fraction of correct instances among all the detected instances, it is defined by

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

Finally, the last metric we interested in is a measure of accuracy, a metric that evaluates the overall quality of the model. The standard accuracy may not be the best metric to look at, as the classes in the present case are highly imbalanced. We will use instead another measure of accuracy, the more appropriated F1-Score, which is the harmonic average of the precision and recall

$$F1 = \frac{2 \times Recall \times Precision}{Recall + Precision}$$

In the literature presented in Section II, sensitivity and specificity are mainly used, as well as the false alarm rate. It is obviously an important metric in the case of seizure detection, we will also communicate results in this sense for comparison. Note that sensitivity and specificity are computed using the Any Overlap Method (OVL), explained in [24]. OVL is more a method to calculate metrics in a permissive way than a metric per se. OVL is based on an event decomposition of the EEG signal. Instead of looking at each segment (also called epoch), OVL only considers the events (i.e. continuous sequence of seizure or background)

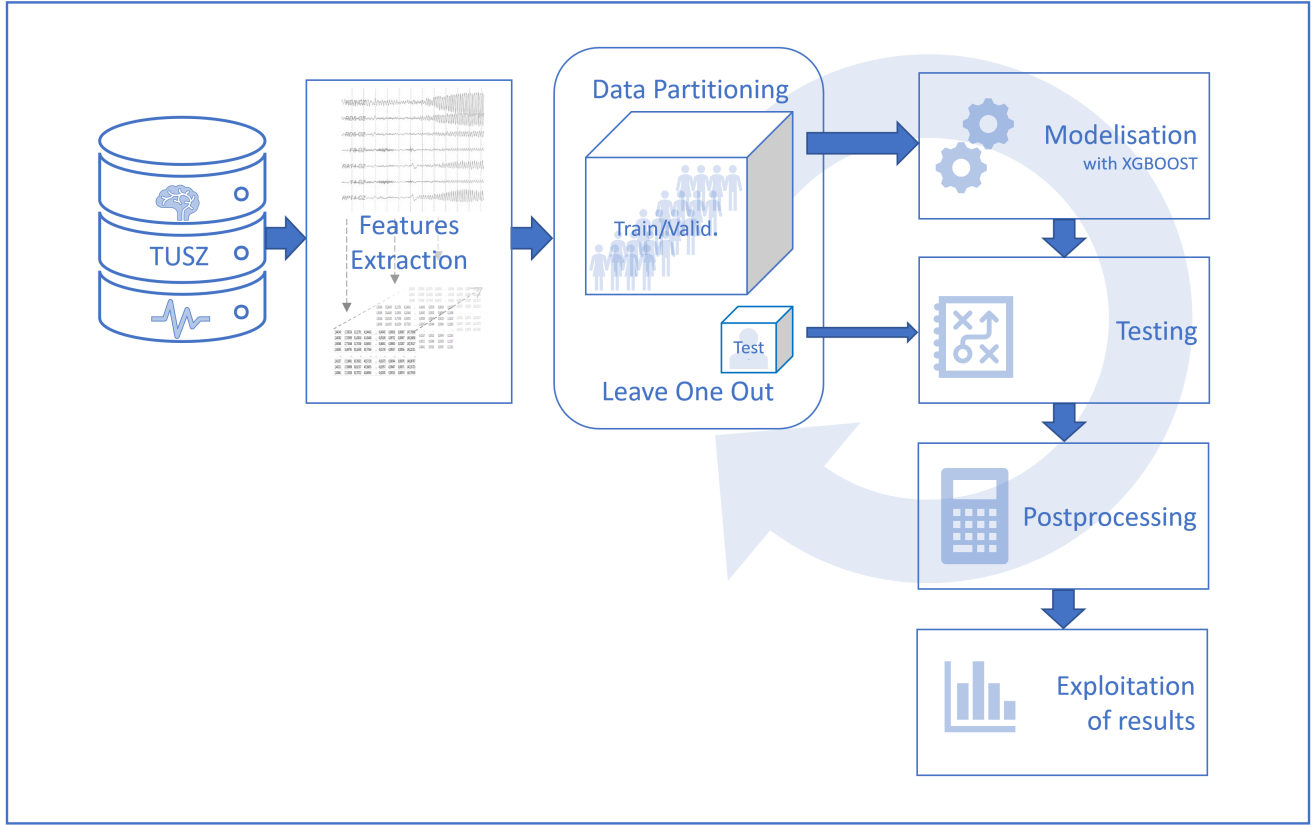


Fig. 3: Second Method based on a LOOCT approach

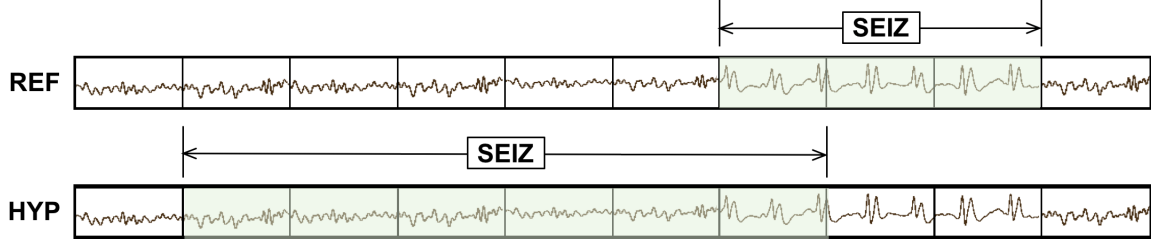


Fig. 4: Illustration of Any Overlap method [24]

that occur. Fig. 4 shows a potential drift of the method. The signal above shows the reference, divided in segments of equal length while the signal below is the hypothesis, i.e. the sequence predicted by the model. While one would count true/false negatives/positives by matching each pair of corresponding segments, OVLP considers that the true positive rate is 100% as soon as the hypothesis overlaps the reference for at least  $\lambda$  segments, where  $\lambda$  is often set to 1 to guarantee that the events do overlap. Thus, in the example, true positive rate is 100% while traditional metrics would only count one segment as a true positive.

In table I, we have deliberately retrieved the distribution of types of seizures in terms of seizure segments instead of events, because it seems to us that the impact on the models of the number of segments is more relevant and reliable. Effectively the global ratios of events and of segments are

different. For instance in AR train, there are 77 events of non specific generalized seizures for 6743 segments. In the LE train, there are 35 events for 10230 segments. It means that the non-specific generalized seizures are way longer in the montage LE. This information is hidden if we count by events.

#### IV. RESULTS AND DISCUSSION

##### Global results

Firstly, in order to get significant results in reasonable time, we started by testing the two described methods on the AR test set. The results are thus corresponding to 36 patients, 899 recordings containing 562 seizure events and more than 49000 seizure segments. The population of each subset is recalled in table III. The events and segments are divided between all type of focal and all type of generalized

|   | AR Train | LE Train | AR Test    | LE Test    | 1st meth. | 2nd meth. | Sens.  | Prec.  | Spec.  | F1     |
|---|----------|----------|------------|------------|-----------|-----------|--------|--------|--------|--------|
| 1 | train    |          | test       |            | x         |           | 35,80% | 31,65% | 92,46% | 33,60% |
| 2 | train    | train    | test       |            | x         |           | 41,74% | 35,17% | 92,50% | 38,17% |
| 3 | train    |          | train-test |            |           | x         | 52,66% | 27,55% | 86,50% | 36,17% |
| 4 | train    | train    | train-test |            |           | x         | 53,87% | 27,57% | 86,20% | 36,47% |
| 5 | train    |          |            | test       | x         |           | 71,65% | 28,67% | 73,04% | 40,95% |
| 6 | train    | train    |            | test       | x         |           | 66,70% | 34,06% | 80,47% | 45,09% |
| 7 | train    | train    | train      | test       | x         |           | 78,72% | 33,54% | 76,41% | 47,04% |
| 8 | train    | train    |            | train-test |           | x         | 72,59% | 39,61% | 83,26% | 51,25% |
| 9 | train    | train    | train      | train-test |           | x         | 71,61% | 40,01% | 83,76% | 51,33% |

TABLE II: List of experiments - selected sets, method and results

| No. of   | patients | recordings | focal events | generalized events | focal segments | generalized segments |
|----------|----------|------------|--------------|--------------------|----------------|----------------------|
| AR Train | 124      | 1220       | 376          | 102                | 22966          | 7976                 |
| LE Train | 129      | 322        | 310          | 100                | 16164          | 11647                |
| AR Test  | 36       | 899        | 312          | 250                | 31779          | 17596                |
| LE Test  | 24       | 49         | 38           | 44                 | 4202           | 2804                 |

TABLE III: Cardinalities of the selected sets

seizures. As we mentioned before, the table I presented a more detailed classification, but the two non-specific sets are dominant and the other types are not enough populated, so for a more refined analysis regarding the results and the content of the corresponding training sets, we have decided to consider the two main seizure classes: focal and generalized seizures.

The results of the experiments are summarized in table II. The ones concerning the AR Test are in the first section and numbered from 1 to 4.

The first observation answers to the question of this work: by increasing the amount of data, we improve significantly the detection score. However, by looking at the F1-score, this needs to be discussed as it is mainly true for the first method using the classic split of the dataset. As one can see, in the experiment 2, by adding LE Train to the training set, all metrics are improved compared to the experiment 1; the F1-score increases by 5% to reach 38%. That is less clear with the second method LOOCT, the sensitivities are higher than with the first method, but we observed a high number of False Positive, which degrades the precision and the specificity. This translates into a stagnant F1-score at 36% for the experiments 3 and 4.

We were expecting a bit more from the second method, that's why we started a second round of experiments with LE Test as test set. This subset is smaller than the AR Test; 82 events against 562 for the AR Test. Note that these subsets have 10 patients in common, these ones were thus filtered out from the training set for the cases where the AR Test set was used in this role. The experiments are numbered from 5 to 9. The contribution of adding the LE train set is confirmed. The use of the second method shows in this case a good improvement of the F1-score, we are 4% over the best F1-Score obtained with the first method (experiment 7). We can

also observe that the addition of the AR Test in the training set do not bring the strong improvement that we could expect. Even, in the seventh experiment, we observe a decrease of the precision and of the specificity as we had in the first round of experiments. This tends to say that the problem is not the second method but the use of the AR Test set as training material. This observation has to be qualified by the difference in size between the AR Test and LE Test sets. Nevertheless, following a discussion with the TUH Team, we learned that they have visually observed that LE files were much cleaner compared to AR files in terms of signal noise ratio and that it would make sense to get bad performances on AR compared to LE files. Moreover, in [20], the importance of the montage and of the reference point is discussed as it practically changes the nature of the waveforms even by using the conversion to the bipolar montage. Another explanation to the problem of adding AR Test in the training set could lie in the properties of the inter-ictal seizure signals. The Fig. 6 shows a recording in the AR Test set. Our system, based on frames of 1 second, detects a seizure state from the beginning of the record. The TUSZ annotations indicates the beginning of the seizure at 7 seconds. The EEG patterns are the same but these "Burst-Suppression" events are not considered as seizure states, as the duration is under 10 seconds. We can imagine the problem of adding these "background" annotated segments in a training set. The intuition is that the false positive rate would be penalised. On the other hand, it is also probably the case in the other subsets dedicated to the training set, the only difference is that the proportion of patients with epilepsy is higher in the AR Test. These observations also explain why performances are so low on the TUSZ dataset. Another important reason to remember is that most TUSZ recordings without seizures are also classified as abnormal, as the EEG recordings are



made in a hospital on patients suffering from neurological problems. Anyway, this advocates for two things: firstly, a better strategy at the level of the model for handling these inter-ictal anomalies; secondly the training data should be selected cautiously.

Comparatively to the performances reported in the TUH literature, these results are of the same order as the ones of their best model [9], [10]. The sensitivity is higher, the specificity a bit lower. In terms of false alarm rate for the patients without epilepsy in the test set, 5 events in 3 recordings were detected on a total of 82 recordings in the second experiment; the false alarm rate is then estimated around 8 FA/24h. If all the recordings without seizure are considered, the false alarm rate rises to 18 FA/24h. This denotes a good capacity of the model to distinguish patients with epilepsy from the others. More, if we compute the scores by records and not by segments, the precision rises to 64,33%, which means that there's a majority of recordings correctly classified on all the recordings detected with a pathology of epilepsy. The other metrics are unchanged.

#### *Results by type of seizure*

If we look at the sensitivities by type of seizure for the first round of experiments in the table IV, we can make the same observations in terms of improvement in function of the used training sets. Moreover, by making a proportional average of these scores with the amount of focal and generalized segments in table III, we can find back the sensitivities in the table II. The main information that we can draw from this view is that the generalized seizures are better detected than the focal ones. This was expected; nevertheless, we have to point out that the proportion in the training set is in disfavour of the generalized seizures, this one oscillating between 25% and 33%. The difficulty to detect focal seizures is due to the fact that they begin in one specific area of the brain with later a potential generalization, while the electrical discharge of the generalized seizures is directly widespread in both sides of the brain.

| Experiment  | 1      | 2      | 3      | 4      |
|-------------|--------|--------|--------|--------|
| focals      | 23,02% | 27,68% | 35,02% | 37,61% |
| generalized | 59,50% | 66,24% | 81,51% | 81,10% |

TABLE IV: Sensitivity by type of seizure

This would advocate to analyse the seizure detection problem with a type-specific approach instead of a generic modeling; however, as argument for the generic approach, 11 patients on the 36 of the AR Test encountered both focal or generalized seizures in the same recordings or in several. Nevertheless, some strategies should be adopted to deal with the different locations of the onsets of the focal seizures. For instance, a new metric that takes into account these preoccupations, such as the importance of characterizing the beginning of a seizure, would be meaningful.

#### *Channel and feature importance*

The idea of type-specific analysis is recalled by the study of the most important features and channels that were the more decisive during the training to discriminate seizures from background. Indeed, the coverage of specific focal seizures in the training set will have an incidence on the importance of the channels. Among all the experiments and the training sets, a feature on two channels emerged, namely the power spectral density on the beta frequency band for the channels P4-O2 and P3-O1. These two particular channels are located symmetrically in the posterior region of the head between the parietal and occipital area. After discussion of our results with neurologists specialized in epilepsy from the University Hospital of Saint-Luc in Brussels, we hypothesized two possible explanations:

- The first one refers indeed to the type of seizure from the location point-of-view. The locations of these features correspond to the place where the focal seizures of occipital origin start. Generalized seizures can also induce ictal changes in this area due to their widespread nature, they are nevertheless a minority in the training sets, focal seizures are most frequently observed. Considering the general statistics of epilepsy among all populations, occipital seizures have been reported to constitute 8% of the population with epilepsy and are occasionally subject to a secondary generalisation [25]. The dominant type of the focal seizures is generally considered to be the Temporal Lobe Epilepsy (TLE) which represents more than 60% of all focal seizures [26]; this type of seizure rarely spreads to the occipital lobes. Taken together, all these statistics do not directly explain the location and the importance of these channels. Moreover, neurologists reveal that a bias can exist in the statistics of hospitals specialized in epilepsy compared to the whole population with epilepsy, as patients with refractory epilepsy (pharmacoresistant epilepsy) are more likely to be redirected through these institutions. The Temple University Hospital is indeed accredited as a level 4 epilepsy center, which means that they have the professional expertise to provide the highest level medical and surgical treatment for patients with refractory epilepsy [27]. To sort out this question, we would need to determine the exact locations of all the included seizures from the dataset. This information is not explicitly given in the TUSZ, but it could be retrieved by processing the .lbl files presented in Section III, which be the subject of a future work.
- The second potential explanation concerns the general physiological activity in the occipital lobe. When reviewing the EEG, neurologists first consider the occipital region; as a normal background activity is generally observed in these channels, with EEG signals less polluted by artefacts like muscles or eyes movements, it helps to interpret the onset of an eventual seizure. If this empiric method reveals indeed a visual discrim-

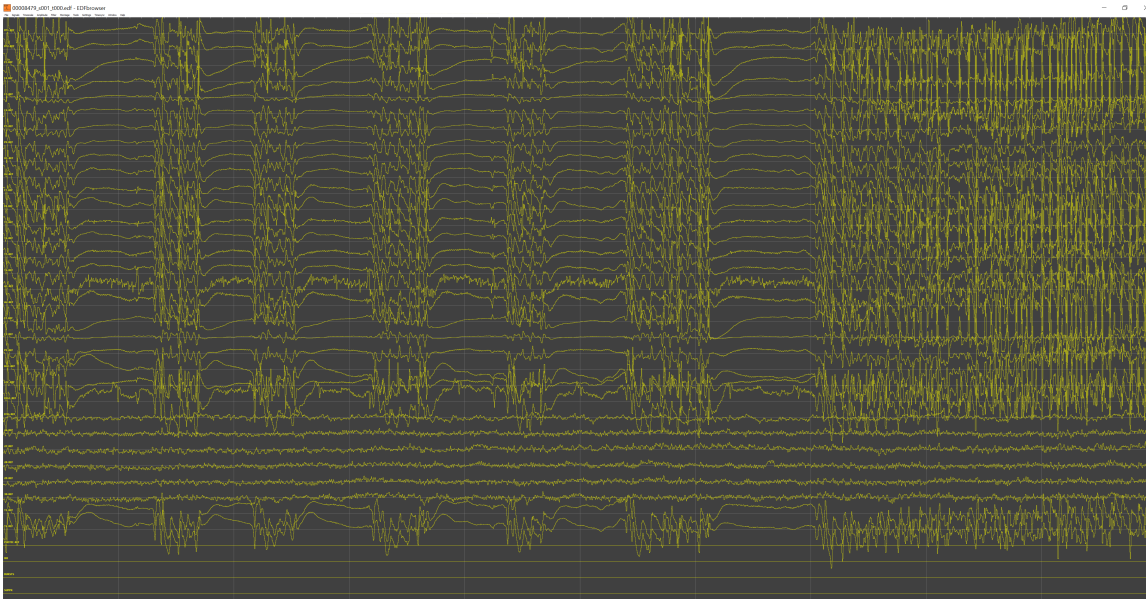


Fig. 6: recording 00008479\_s001\_t000

ination of ictal and interictal phases, it is likely that the algorithm has found the same rule. Furthermore, a study from the TUH aims to measure the performance of their model if fewer channels are considered [28]. The removed channels were chosen by using domain knowledge instead of using a proper automated and optimized selection process. If only two channels were to be preserved, one of them would be located in the occipital region and the other bound to the CZ electrode for its central location. The removal of the other channels is justified either by their susceptibility to noise or artefact interference, or by their less strategic location. So, to conclude this point, experts seem to be concerned by the activity in the posterior region of the head and our interpretable models suggest this location for the most important features. The validation of this relation will also be the subject of a future work.

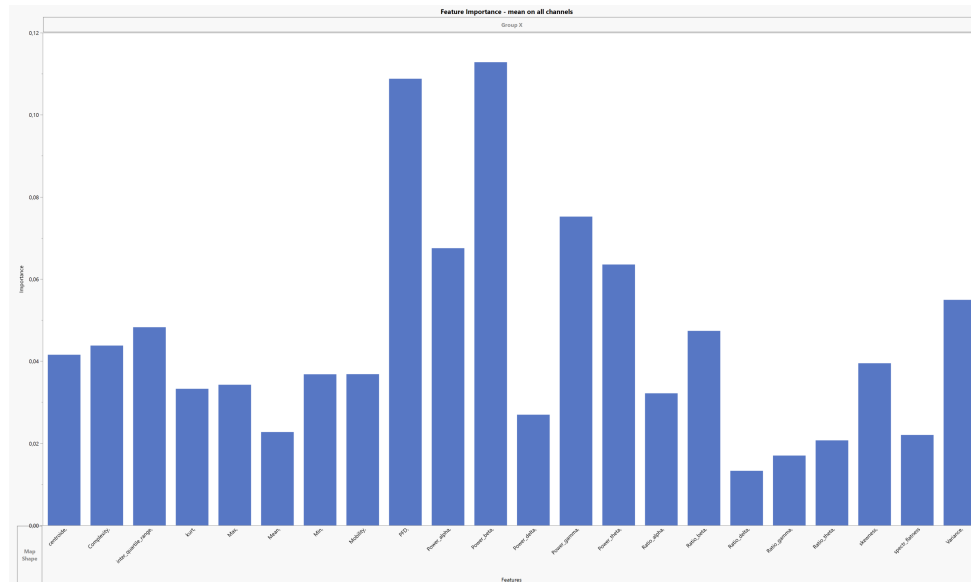
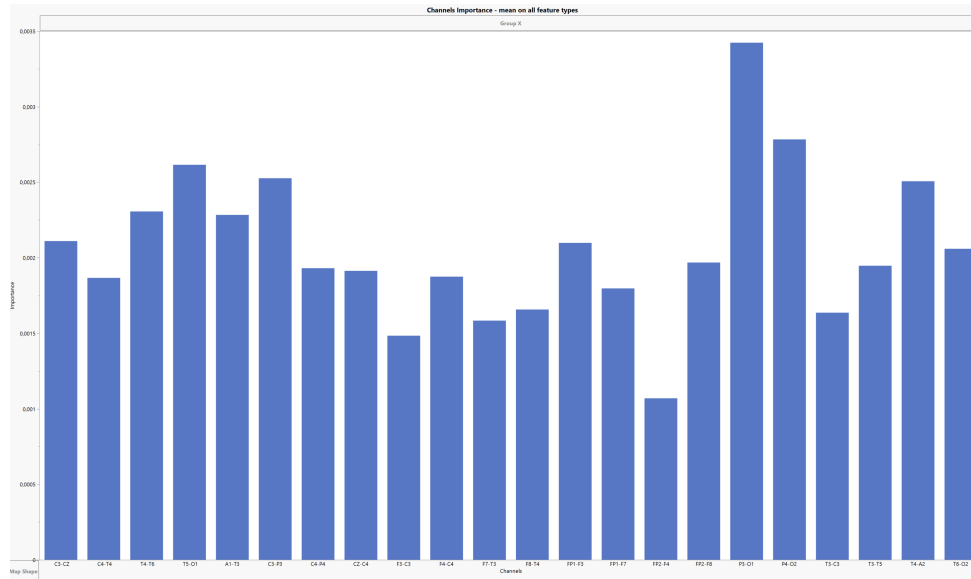
Below the power spectral density on the beta frequency band for the channels P4-O2 and P3-O1, we can find other significant combinations. As there are 22 features for 22 channels, instead of enumerating them, we have done a global analysis presented in two figures. Fig. 7 shows the averaged channel importance over all the features and Fig. 8 shows the averaged feature importance of all the channels. For the channel importance, we find first the two already mentioned channels between the occipital and parietal area, then some others between the central and the parietal area or between the temporal and occipital area or purely temporal. Two things are interesting to note: firstly it's not symmetric, as if the channels in the right hemisphere were more important; secondly the less important channels are related to the frontal area, a sign of the blindness of the EEG in this zone as mentioned in [28]. For the feature importance, the power spectral density on the beta frequency band confirms to be the first on average, but very closely and far above the others, we

find the Petrosian Fractal Dimension. Various studies have shown the interest of using fractals for biosignals analysis and for seizure detection [29], [30]. For the other features, we notice the good performance of the other power spectral densities in alpha, gamma and theta frequency bands and of the variance.

## V. CONCLUSION AND PERSPECTIVES

In this paper, we analysed the problem of seizure detection using machine learning techniques on the Temple University Hospital EEG Seizure Corpus and under the perspective of the increase of the amount of data. By using both a classic machine learning approach and a Leave-One-Out method on two subsets, we generated experiments based on training sets which have the characteristics to be large and heterogeneous, which is not common in the current state-of-the-art. Using an efficient implementation of the gradient boosting algorithm called XGBoost, we showed the potential of improvement that can be done by adding more training data, even if this one should be selected cautiously. We presented results of the same order of magnitude than the corresponding literature based on deep learning architectures. The performances by type of seizure were also analysed and showed that generalized seizures are easier to recognize, even if they are under-represented in the training set. Finally the interpretable property of the XGBoost algorithm shows that a feature linked to two specific EEG channels between the parietal and occipital area is statistically the most discriminant to distinguish a seizure from a non-seizure (considering segments of one second and a generic modeling done with the proportions by type of seizure proposed by the TUSZ).

In further works, we would like to extend this approach to the futur releases of the TUSZ with more data. A larger dataset means that the LOOCT method won't be necessary anymore. Based on the actual findings, we are planning to



study an architecture that would be more able to discriminate ictal phase from inter-ictal, and also a type of seizure from another, either by using a multiclass classification approach by distinguishing focals and generalized seizures, either by leaving aside the generic approach and by focusing on type-dependent models. The use of interpretable models is still relevant and it would be interesting to see how the importance of the features evolve in function of the focal seizures and their onset locations. We are also interested to do some process mining based on the by-channel annotations by looking at the order of (dis)appearance of a seizure on particular channels in order to identify channels that are the most critical. One of the goals would be to converge to a system capable of early detection. Such a fast detection is critical for clinicians because it allows nurses to react as

quickly as possible with the appropriate actions.

From a technical point of view, deep learning models could be tested again. Studying other sources of signals such as electrocardiograms (ECG) would also be interesting but once again, we suffer from a lack of data resources.

This work was carried out as part of the BIG DATA PLATFORM project (<https://www.cetic.be/Big-Data-Platform>), supported by the Région Wallonne (convention number 7481).

We would like to express our deep gratitude to Professor Susana Ferrao Santo, neurologist and epileptologist (Refractory Epilepsy Centre, University Hospital of Saint-Luc, Brussels), and Simone Vespa, Research Assistant (Institute of

Neuroscience, Catholic University of Louvain, Brussels), for their interest in our work and for the constructive discussions we had.

We would also like to thank Professor Joseph Picone and his colleagues Vinit Shah and Meysam Golmohammadi (Department of Electrical and Computer Engineering, Temple University, Philadelphia, PA, USA) for their valuable feedbacks to our questions related to the TUH corpus.

## REFERENCES

- [1] C. E. Elger and C. Hoppe, "Diagnostic challenges in epilepsy: seizure under-reporting and seizure detection," *The Lancet Neurology*, vol. 17, no. 3, pp. 279–288, 2018.
- [2] S. J. M. Smith, "Eeg in the diagnosis, classification, and management of patients with epilepsy," *Journal of Neurology, Neurosurgery & Psychiatry*, vol. 76, no. suppl 2, pp. ii2–ii7, 2005.
- [3] A. H. Shueb, *Application of machine learning to epileptic seizure onset detection and treatment*. PhD thesis, Massachusetts Institute of Technology, 2009.
- [4] M. Bell, "The ontogeny of the eeg during infancy and childhood: Implications for cognitive development," pp. 97–111, 1998.
- [5] K. M. Tsiouris, V. C. Pezoulas, M. Zervakis, S. Konitsiotis, D. D. Koutsouris, and D. I. Fotiadis, "A long short-term memory deep learning network for the prediction of epileptic seizures using eeg signals," *Computers in biology and medicine*, vol. 99, pp. 24–37, 2018.
- [6] R. G. Andrzejak, K. Lehnertz, F. Mormann, C. Rieke, P. David, and C. E. Elger, "Indications of nonlinear deterministic and finite-dimensional structures in time series of brain electrical activity: Dependence on recording region and brain state," *Physical Review E*, vol. 64, no. 6, p. 061907, 2001.
- [7] R. Hussein, H. Palangi, R. Ward, and Z. J. Wang, "Epileptic seizure detection: A deep learning approach," *arXiv preprint arXiv:1803.09848*, 2018.
- [8] D. Ahmedt-Aristizabal, C. Fookes, K. Nguyen, and S. Sridharan, "Deep classification of epileptic signals," *arXiv preprint arXiv:1801.03610*, 2018.
- [9] M. Golmohammadi, S. Ziyabari, V. Shah, S. L. de Diego, I. Obeid, and J. Picone, "Deep architectures for automated seizure detection in scalp eegs," *arXiv preprint arXiv:1712.09776*, 2017.
- [10] M. Golmohammadi, S. Ziyabari, V. Shah, E. Von Weltin, C. Campbell, I. Obeid, and J. Picone, "Gated recurrent networks for seizure detection," in *Signal Processing in Medicine and Biology Symposium (SPMB)*, 2017 IEEE, pp. 1–5, IEEE, 2017.
- [11] S. Itani, F. Lecron, and P. Fortemps, "Specifics of medical data mining for diagnosis aid: A survey," *Expert Systems with Applications*, 2018.
- [12] M. V. A. T. . P. J. Ferrell, S., "The temple university hospital eeg corpus: Electrode location and channel labels," 2018.
- [13] T. Chen and C. Guestrin, "Xgboost: A scalable tree boosting system," in *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, pp. 785–794, ACM, 2016.
- [14] Y. Zhang, W. Zhou, S. Yuan, and Q. Yuan, "Seizure detection method based on fractal dimension and gradient boosting," *Epilepsy & Behavior*, vol. 43, pp. 30–38, 2015.
- [15] G. Wang, Z. Deng, and K.-S. Choi, "Detection of epileptic seizures in eeg signals with rule-based interpretation by random forest approach," in *International Conference on Intelligent Computing*, pp. 738–744, Springer, 2015.
- [16] Y. Jiang, Z. Deng, F.-L. Chung, G. Wang, P. Qian, K.-S. Choi, and S. Wang, "Recognition of epileptic eeg signals using a novel multiview task fuzzy system," *IEEE Transactions on Fuzzy Systems*, vol. 25, no. 1, pp. 3–20, 2017.
- [17] M. Golmohammadi, A. H. H. N. Torbati, S. L. de Diego, I. Obeid, and J. Picone, "Automatic analysis of eegs using big data and hybrid deep learning architectures," *arXiv preprint arXiv:1712.09771*, 2017.
- [18] V. Shah, E. von Weltin, S. Lopez, J. R. McHugh, L. Veloso, M. Golmohammadi, I. Obeid, and J. Picone, "The temple university hospital seizure detection corpus," *arXiv preprint arXiv:1801.08085*, 2018.
- [19] I. Obeid and J. Picone, "The temple university hospital eeg data corpus," *Frontiers in neuroscience*, vol. 10, p. 196, 2016.
- [20] S. Lopez, A. Gross, S. Yang, M. Golmohammadi, I. Obeid, and J. Picone, "An analysis of two common reference points for eegs," pp. 1–5, 2016.
- [21] A. T. Berg, S. F. Berkovic, M. J. Brodie, J. Buchhalter, J. H. Cross, W. van Emde Boas, J. Engel, J. French, T. A. Glauser, G. W. Mathern, et al., "Revised terminology and concepts for organization of seizures and epilepsies: report of the ilae commission on classification and terminology, 2005–2009," *Epilepsia*, vol. 51, no. 4, pp. 676–685, 2010.
- [22] F. S. Bao, X. Liu, and C. Zhang, "Pyeeeg: an open source python module for eeg/meg feature extraction," *Computational intelligence and neuroscience*, 2011.
- [23] R. Miller, "Theory of the normal waking eeg: from single neurones to waveforms in the alpha, beta and gamma frequency ranges," *International journal of psychophysiology*, vol. 64, no. 1, pp. 18–23, 2007.
- [24] S. Ziyabari, V. Shah, M. Golmohammadi, I. Obeid, and J. Picone, "Objective evaluation metrics for automatic classification of eeg events," *arXiv preprint arXiv:1712.10107*, 2017.
- [25] J. Duncan, "Chapter 15 - occipital and parietal lobe epilepsies." ILAE lecturer, fifteenth teaching weekend on epilepsy, [https://www.epilepsysociety.org.uk/lecture-notes-0#.XHj\\_6YhKguU](https://www.epilepsysociety.org.uk/lecture-notes-0#.XHj_6YhKguU), 2015. [Online; accessed 01-march-2019].
- [26] B. Diehl and J. Duncan, "Chapter 13 - temporal lobe epilepsy." ILAE lecturer, fifteenth teaching weekend on epilepsy, [https://www.epilepsysociety.org.uk/lecture-notes-0#.XHj\\_6YhKguU](https://www.epilepsysociety.org.uk/lecture-notes-0#.XHj_6YhKguU), 2015. [Online; accessed 01-march-2019].
- [27] <https://www.templehealth.org/services/neurosciences/patient-care/programs/epilepsy>. [Online; accessed 01-march-2019].
- [28] V. Shah, M. Golmohammadi, S. Ziyabari, E. Von Weltin, I. Obeid, and J. Picone, "Optimizing channel selection for seizure detection," in *2017 IEEE Signal Processing in Medicine and Biology Symposium (SPMB)*, pp. 1–5, IEEE, 2017.
- [29] T. Q. D. Khoa, V. Q. Ha, and V. V. Toi, "Higuchi fractal properties of onset epilepsy electroencephalogram," *Computational and mathematical methods in medicine*, vol. 2012, 2012.
- [30] Y. Zhang, S. Yang, Y. Liu, Y. Zhang, B. Han, and F. Zhou, "Integration of 24 feature types to accurately detect and predict seizures using scalp eeg signals," *Sensors*, vol. 18, no. 5, p. 1372, 2018.