



EEG-based Epileptic Seizure Detection Using DconvNET

Suresh Nalla^{1,*}, Seetharam Khetavath²

¹Research Scholar, Department of ECE, Chaitanya Deemed to be University, Kishanpura Hanamkonda, Warangal, Telangana, India

²Professor, Department of ECE, Chaitanya Deemed to be University, Kishanpura, Hanamkonda, Warangal, Telangana, India

Emails: sureshnalla44@gmail.com; seetharamkhetavath@gmail.com

Abstract

Epilepsy is a neural condition that is rather prevalent and affects a sizeable portion of the average population all over the world. Throughout its history, the illness has constantly be located of significant status in the pitch of biomedicine due to the dangers it poses to people's health. Electroencephalogram (EEG) recordings are a method that may be utilized to evaluate epilepsy, which is defined by the occurrence of seizures that occur repeatedly and without any apparent cause. Electroencephalography, often known as EEG, is a method that is utilized to assess the electric movement located within the brain. The examination of electroencephalogram data is an essential component in the field of epilepsy research, since it allows for the early detection of epileptic episodes. On the other hand, the generation of models that are independent of individual characteristics is a significant challenge. Extensive efforts have been directed to the creation of classifiers that are tailored to specific patients. In this thesis, the cross-patient viewpoint is the primary focus of investigation; nevertheless, the heterogeneity of EEG patterns among people presents a challenge to this investigation. An examination of the similarities and differences of the pattern recognition algorithms that are applied for the diagnosis of epileptic episodes based on EEG data was taken. SVM (Support Vector Machine) and KNN (K-Nearest Neighbor) were the approaches that were under consideration for evaluation. According to the findings of our analysis, the two approaches exhibit comparable levels of performance; however, KNN attained a slightly greater level of accuracy in some situations on occasion.

Keywords: Epilepsy, EEG, SVM, KNN, Convolutional neural networks

1. Introduction

Throughout 70 million individuals throughout the world are affected with epilepsy, which is a neurological illness that is rather prevalent. Around 2.4 million cases of the illness are reported each year, and 85 percent of the affected population is located in countries that are still in the process of developing. According to estimates, around fifty percent of epilepsy cases begin during childhood or adolescence, and the condition can present itself in individuals of any age. Due to the fact that individuals who are affected by epilepsy have a two to threefold greater risk of early mortality in comparison to those who are not affected by the ailment, the examination of epilepsy is a very imperative element in the ground of biomedical investigate. Epilepsy is defined by the repeated and spontaneous occurrence of seizures, which can be classed as either partial, originating from a specific region of the brain, or comprehensive, distressing the entire brain. Epilepsy is a neurological condition that requires medical attention. When epileptic movement initiates in one cerebral hemisphere of the brain, this is referred to as a partial starting. On the other hand, a generalized onset occurs when both hemispheres of the brain are simultaneously engaged. The specific area of the brain from which epileptic activity originates and the degree to which it spreads are two factors that determine the sort of seizure that a person experiences. The seizures range from mild attentional impairments to severe convulsions that last for a lengthy period of time. Our current understanding of the characteristics of an epileptic brain is limited, despite the fact that these characteristics are extraordinarily complex. Epilepsy can momentarily express itself in a variety of ways, including moments of altered awareness, modest motor skill abnormalities, muscular spasms, and a wide range of other symptoms. In addition, epileptic seizures frequently occur on their own without any interference from outside sources, and as a result, they are sometimes

overlooked. In the case of seizures, ongoing work is expected; this method presents a significant amount of technological challenge. An electroencephalogram, popularly called an EEG, is a procedure used to measure the electric movement in the brain. This technique may be used to assess epilepsy. There has been noted a trend toward increasing interest in the interpretation of the states of the brain in the form of EEG records. In these recordings, the electrical activity of a slice of neural tissue does not signify its function but rather gives only an input-output relation, a black box of causes and effects impulses that were created by the brain through electrodes that were attached to the subject's skull are documented. This is followed by analysis by a computer which dissects these impulses. The actual identification of epilepsy by using EEG data and especially in case of detailed recordings may be a rather a time-consuming process [5]. The challenge of characterizing and analysing the EEG signal arises from its very non-linear and non-stationary nature. However I suppose that is a therapy that is incredibly well known and very cheap. Another critical area of epilepsy is the usage of electroencephalogram (EEG) data for the differentiation of seizures at their preliminary stages. In cases of a simple focal seizing it is crucial that it is diagnosed early enough so that the neurostimulation can be given and therefore prevent it from generalize to extra parts of the brain. Therefore, there is a need for discovery an optimal technique for the task of automatic recognition of seizures [6]. Ever since, there has been an effort to progress a processor assisted scheme for diagnosis of epilepsy so as to improve the management of those who are affected. This is important because EEG recording processing time is greatly reduced with the use of automation and more people can be treated with the technology [7]. A lot of effort has been complete on classifiers which have been created for specific patients only. The second main difficulty is that creating models non-tied to concrete patients is less feasible for a number of reasons because the inter-subject variability in the spectrum of EEG is large [8]. Alas, until now, patient-independent classifiers have not reached the same levels of performance as patient-specific classifiers, even though classifiers constructed with the use of patients' data have been proven effective [9]. Across a wide range of participants, the determination of this investigation is to uncover electroencephalogram outlines that are connected with epileptic seizures. In the event that classifiers demonstrate enough accuracy, independent of the specific patients involved, substantial processing resources might be saved. The utilization of this method would result in a more cost-effective inclusion into automated detection systems. There have been a number of studies that have utilized a variety of ML methods to identify seizures in EEG data [9-14]. Despite this, the majority of research has focused on classifiers that are tailored to specific patients [9]. The SVM and the KNN algorithms were compared in a prior work [13], which resulted in improved accuracy in a patient-specific setting. As remote as we are alert, there have been no earlier trainings in the field of patient-independent research that have explicitly compared and contrasted these two methods. This research investigates the SVM and KNN algorithms for identifying epileptic episodes by utilizing EEG data and compares and contrasts their respective capabilities. Because of the unpredictability in EEG data, the primary focus of this work will be on investigating patient-independent classifiers, which are more complicated than other types of classifiers (9). A seizure is the predominant symptom and clinical concern of a disease that is characteristic of epilepsy that is continuous. In the course of an epileptic seizure, it is sometimes difficult to report on the clinical data on the patient's subjective perception as well as objective findings. This is so, because there are so many options, through which a particular symptom could manifest itself. These include the origin place of the brain, snooze awaken series and level of development of the brain and several others, all which the potential of a seizure occurrence [16]. IOT domain also contributes to identifying the elliptic seizure detection. It uses IOT devices through which the doctors and patients are sent the secure data.

These parameters of epileptic episodes may be examined by the analysis of recorded electro encephalogram (EEG) data. While seizures are occurring, recordings of bioelectric activity in the brain that are collected are called EEG and have their own elder patterns from normal, non-epilepsy individuals. Consequently, electroencephalogram analysis helps to distinguish between data concerning epileptic seizures and data concerning non-epileptic seizures, as well as between the phases of a seizure. The rampant characteristics of Epilepsy mean that predictions can be made of an oncoming seizure by noting a change in the Electroencephalogram (EEG) patterns preceding a seizure. To successfully implement this prediction framework, there is a need for a robust automated system to identify normal, pre-ictal, and intermittent manners. This baseline part includes the EEG information of a vigorous subject, the pre-ictal phase is dedicated to the variations of EEG before a seizure, and the inter-ictal is to show the variations that occur throughout a seizure event. Concerning this recognition approach, there are two crucial aspects that need to be evaluated. First and foremost, the selection of features to excerpt from the EEG information, and more specifically, the approaches that are utilized for feature extraction, is the most important component. The third phase entails the selection of analytical methods for the application to the acquired data for the purpose of stage identification, with a particular emphasis on classification approaches [2].

In scalp electroencephalography, a station is created by the voltage alteration that exists amongst two conductors that have been identified. The human brain is capable of producing five distinct patterns of brain waves, each of which is associated with a certain frequency range. The information that these brain waves transmit about the state of the brain is extremely important. Delta, theta, alpha, beta, and gamma are the occurrence bands that contain the

frequencies. All of these occurrences are associated with various mental states, and an unanticipated disruption in the electric movement of the brain may apparent itself, as demonstrated by an epileptic seizure. In order to recognize these indicators, one can moreover focus on a small number of channels that meet particular requirements or investigate all of the frequencies that are available [17]. The enhancement of seizure detection by the decrease of channels is a significant improvement. When performing an EEG it is quite common to employ many electrodes so as to obtain records from diverse sites of the head. This is done so that all the data that is to be transcribed has to be transcribed in a correct manner. In cases where there are additional channels, such as in the current study, aggressive use of the number of channels can lead to overfitting. This is particularly so where there are numerous channels, for it becomes merry go round of power. By using channel selection approach, it is possible to decrease the dimensionality of the feature design and hence reduce the resources required to undertake feature extraction and classification. In relations of the number of electrodes data, the study conducted in 2015 showed that it is possible to have fewer ELECTROENCEPHALOGRAM channels, which ranges between 10 to 30 percent to be efficient in classification tasks without straining performance. This was realised when the researchers were undertaking their surveys. Some of the signal features like vary and entropy can help to identify channels for specific aim like detecting seizures in eeg patterns based on some research evidence. [17]. This has been proved through the use of research. In the present work, a seizure detection system was developed using a SOM neural network for the purpose of an independent research study. This system was meant to detect seizures and yet it misdiagnoses me for 8 minutes. In all 24 of the prolonged EEG recordings underwent detector independent seizure identification. To classify seizures from the 8 channel from the 18 channel scalp EEGs the neural network was trained using 98 samples. This was done in order to achieve the stated aim. For the processing of the respective EEG epochs wavelet processing was done. An accuracy rate of one hundred percent was achieved again because the system can accurately identify 56 of the 62 possible seizure cases offered to the algorithm. As mentioned in reference [12], it was observed that the false positive rate of the mentioned algorithm was 0.71 ± 0.79 hours per hour. Thus, with the view of building a classifier to perform a personalized seizure detection, a research was conducted. As a result of the investigation, it is possible to establish that the used classification strategies contain two types: KNN and SVM. When the algorithms were benchmarked for their performance on 10 different cases it was found out that both the algorithms had a similar overall percentage of response accuracy of 80 percent but KNN was slightly superior to SVM in actual performance. Besides this, they proved that the onset sensitivity is hundred percent perfect [13].

In this particular work, an approach that may be employed in developing models that are capable of diagnosing epileptic episodes is first introduced and then analysed. The study employed scalp electroencephalogram (EEG) data and in the analysis implemented a patient-tailored approach. When applying the method, the 173 test seizures could be identified with net accuracy of 96 percent, obtained through several seizures in each of the 24 patients and analysed with the help of 916 hours of EEG data. A Support Vector Machine was applied in order to classify the feature vector, and the CHB-MIT database was utilized in the research [9][15]. The research that they conducted utilized the identical dataset that we did, with just a few minor adjustments that were mentioned in section 3.3. With the use of an offline detector known as the Reveal Algorithm [14], which makes use of a neural network, they compared their findings with those of the revealed algorithm. The neural network was trained for a much bigger number of epochs by making use of a larger patient sample with the intention of improving its performance. Consequently, in contrast to the classifiers described in [9], it displayed competence in the management of samples that were not specifically customized for any one particular individual.

2. Methodology

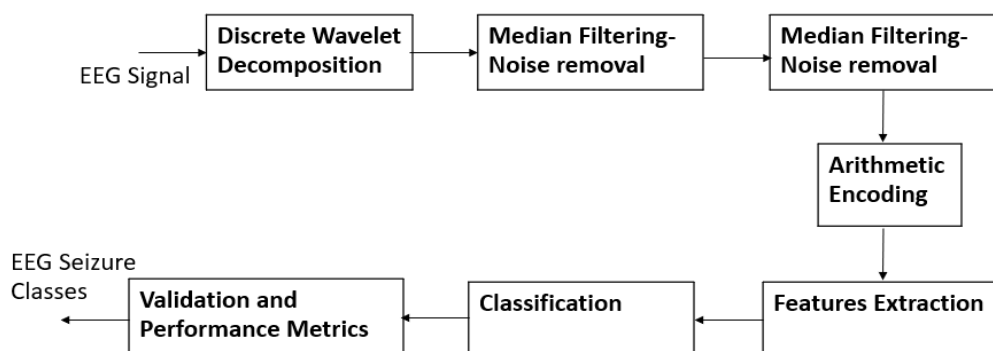


Figure 1. Proposed Methodology Diagram

a. Discrete Wavelet Transforms

Continuous wavelet transforms (CWT) are often characterized as a nonredundant sampling variation of discrete wavelet transforms (DWT), which are also known as discrete wavelet transforms. In order to describe a discrete period sequence, $x(n)$, as a collection of wavelet constants, the Discrete Wavelet Transform (DWT) does its best to achieve this goal. Derivation of the coefficients is accomplished by the utilization of a continuous wavelet transform, which is typically designed to provide an orthogonal or biorthogonal basis function set. There are many different formulations of DWT, and each one has its own set of characteristics. This particular section is only concerned with orthogonal wavelets that have finite support. In order to ensure that there is no repetition in the representation, orthogonal bases are utilized. The use of orthogonal representations makes direct approaches necessary for reconstruction and deconstruction more accessible. In many cases, efficient DWT techniques need fewer CPU resources than those required for a rapid Fourier transform. The DWT may be interpreted from a variety of perspectives that are all equal. Within the context of a filter bank design, the DWT is the subject of this discussion. Specifically, this article presents two FIR filters, each of which is characterized by a total of L coefficients for individual filters. The first is titled high pass filter while the second one is titled low pass filter; the two filters are activated and deactivated at a frequency of half the sampling rate. The use of such filters in an iterative manner makes it possible for this approach, known as the Discrete Wavelet Transform (DWT), to be unique in some way. From the beginning, the filters are pragmatic to the time sequence that is being input in order to generate low-pass and high-pass mechanisms, which are denoted by $x_1(n)$ and $x_2(n)$, respectively:

$$\begin{aligned} x_1(n) &= \sum_{k=0}^{L-1} c_k x(n-k) \\ x_2(n) &= \sum_{k=0}^{L-1} d_k x(n-k) \end{aligned}$$

where c_k are constants of the low-pass filters, and d_k – coefficients of the high-pass ones. To establish the link between the two groups of strainer coefficients it is usually expected that the high-pass filter is to be derived from the low-pass filter. This is normally done by the discontinuous flip approach.

$$d_k = (-1)^k c_{L-k}$$

Since $x_1(n)$ covers the lower frequency and $x_2(n)$ covers higher frequency band, the output of both filters is band limited and has bandwidth equal to one half that of the input sequence entered. The information contained within these two-time series is redundant as a consequence of the fact that the outputs of each filter have a bandwidth that is half of the original bandwidth of $x(n)$. When the outputs of the two filters are subsampled, it is possible to get a sampling rate that is fifty percent of the rate that was initially determined. Subsampling is achieved by removing every other sample from inside each sequence through the elimination process. The fact that both signals make use of the whole bandwidth is demonstrated by this resampling. Aliasing occurs in both of the components during the resampling process as a consequence of the poor strategy of the two FIR sifters respectively. It is feasible to engineer the two filters in such a manner that their aliasing effects cancel each other out, which will minimize the artifacts that are caused by aliasing. This will allow for a reduction in the amount of artifacts that are produced. Using the two subsampled sequences as their starting point, it is feasible to reproduce the original signal after filters have been constructed to eliminate aliasing. This is achievable since the filters have been implemented. To reconstruct the inventive input categorization from the two subsampled categorizations, you must first up sample the two subsampled indicators, then apply filtering, and then combine the two mechanisms. This will allow you to reconstruct the original input sequence. It will be possible for you to reproduce the first sequence of inputs thanks to this. The synthesis filters and the decomposition filters are identical in an orthogonal wavelet transform that makes use of FIR filters. The only difference between the two is that the impulse responses of the synthesis filters and the decomposition filters are presented in the opposite direction of time. As shown in Figure 7A, the procedure of disintegrating a indication into two subsampled categorizations and subsequently rebuilding the original signal is depicted. An illustration of this technique may be found in the diagram. In Figure 7, the symbols $\uparrow 2$ and $\downarrow 2$ are used to illustrate the processes of up sampling and down sampling, respectively. The transmission purposes of the low-pass and high-pass disintegration filters are represented by the symbols $H_1(z)$ and $H_2(z)$, respectively. These symbols are assigned to the filters. The numbers $H_3(z)$ and $H_4(z)$ and respectively are used to represent the transfer functions of the synthesis filters that relate to each other.

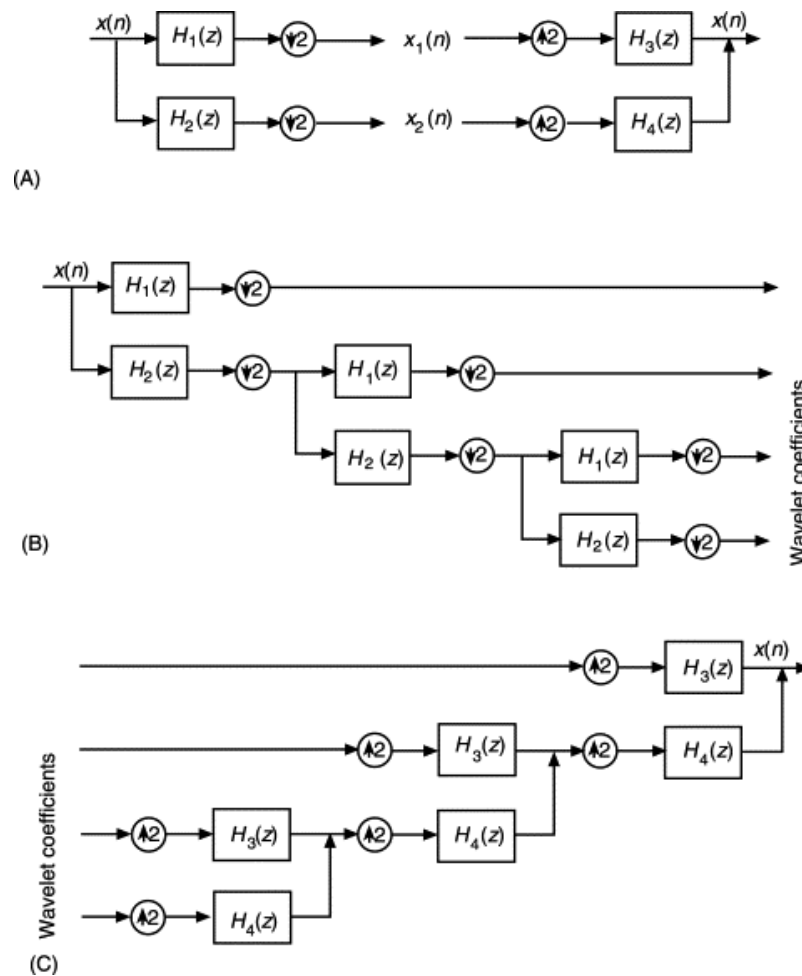


Figure 2. Wavelet transform with filter banks

3. Median Filter

Through the process of signal processing, it is usually beneficial to reduce the amount of noise present in a picture or signal. Removing noise is a common application of the median filter, which is a nonlinear digital filtering algorithm. To advance the outcomes of advanced analysis, such as edge detection in pictures, noise reduction is a pre-processing procedure that is routinely used. Because it maintains the edges of the picture while successfully removing noise in certain situations, median filtering is widely used in digital image processing technology.

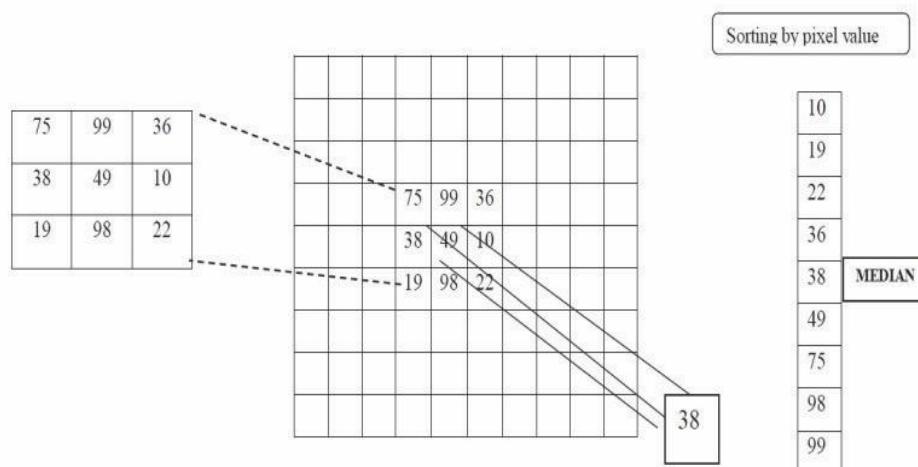


Figure 3. Median filter functioning

The underlying idea behind the median filter is that iterative processing of the signal is performed, with each input being replaced with the median of the values that are close to it. The arrangement of neighbours is known to as the "window," and it proceeded in a methodical manner over the whole signal, one entry at a time. The simplest obvious window for one-dimensional signals is made up of the entries that came before and after it. On the other hand, for two-dimensional (or higher-dimensional) signals, such as images, it is possible to create more complex window configurations, such as "box" or "cross" patterns. The median is the value that is firmly recognized as the centre value in a dataset that has an odd number of elements. This is determined by numerically ordering all of the entries.

4. k-nearest neighbors (knn) algorithm

Within the context of determining the arrangement of a new model point, the KNN technique makes use of a dataset that already contains known classifications. Examples are the most efficient means of providing clarity on this topic.

Consider the following scenario: a financial organization maintains a database that contains information about individuals, including their credit scores. It is quite probable that these aspects are connected to the individual's financial characteristics, such as their salary, whether or not they are homeowners, and other relevant criteria, which would be employed to evaluate the individual's credit rating. The evaluation of a person's credit rating, on the other hand, presents a significant financial burden, which is why the bank has taken steps to reduce costs. They are aware of the fact that, according to the intrinsic qualities of credit ratings, individuals who possess the same financial information would acquire the same credit ratings. They want to make use of the data that is already available in order to make a prediction regarding the credit rating of a new client without carrying out all of the computations. A botanist has the intention of conducting research on the floral variety that may be found within a large meadow. Despite this, he is unable to devote sufficient time to the comprehensive study of each bloom, and he does not possess the financial resources necessary to hire others with the necessary expertise to assist him. He gives them instructions to quantify a number of specific characteristics of the flowers, including the size of the stamens, the number of petals, the height, the colors, and the ratios of the flower heads, and then to enter this information into a computer. Then, he gives the computer instructions to examine them in comparison to a database of samples that has already been identified, and he forecasts the variety of each bloom based on the characteristics of the samples. In most cases, we start with a dataset in which every single data point is classified into a certain category. We want to be able to make a prediction about the categorization of a new information point constructed on the categories that have been defined for the observations that are stored in the database. In addition to giving insights into the characteristics of a wide variety of categories, the database also acts as our training set. The classification problem is the process of classifying a new observation, and there are a number of different approaches that may be used to solve an issue of this nature. When evaluating the categorization of the new observation, we make reference to the classifications of the observations in the database that are the most comparable to the new observation. One of the most difficult problems to solve is determining how similar two observations are to one another. The evaluation of the resemblance amongst two colors is fundamentally dissimilar from the evaluation of the resemblance amongst two sentences that consist of written words. In light of this, it is necessary to create a comparison strategy before attempting to evaluate the degree of resemblance that exists between two specific observations. Considering that our data may include a variety of categories, such as numerical values, colors, geographical locations, or binary replies, the key difficulty is that each of these categories requires a different set of approaches in order to successfully evaluate similarities. It would appear that the primary concern is the structure of the information included in the database in order to make the comparison of observations easier. The transformation of all characteristics into numerical representations is a typical way for accomplishing this goal. For example, colors may be converted into RGB values, locations can be converted into latitude and longitude, and Boolean values can be converted into ones and zeros. We are able to envision a space in which each feature is represented as a different dimension, and the value of each observation serves as its coordinate within that dimension. This is made possible by quantifying all of the variables. Consequently, our observations are transformed into spatial locations, which enables us to evaluate the distance between them as a measure of how similar they are to one another. Once we have developed a system that can evaluate the degree of similarity between two observations, we will need to solve the difficulty of determining which observations from the database are sufficiently comparable to our new observation in order to have an impact on its categorization. This problem may be solved in a number of different ways, including doing an analysis of all the data points that are located within a particular radius of the new sample point or picking a restricted subset of the points that are closest to the new sample point. An analysis of the latter approach using the KNN Algorithm is going to be done at the moment.

KNN also refer to the K-Nearest Neighbours method is the supervised machine learning classification method used for solving classification and regression problems. Evelyn Fix and Joseph Hodges in 1951 when they first came up with the technique and Thomas cover the one who refined it. In light of the information established in this article, the KNN algorithm and its concepts and operations are discussed.

In the ground of DL, KNN is one of the family classification that distinctively characteristic as fundamental and primary one. Pattern recognition, data mining as well as intrusion detection apply themselves heavily within this method which comes close to that of supervised learning.

Unlike, for example, GMM, which assumes Gaussian distribution, this is one of the advantages of the proposed method due to its non-parametric nature, which does not require assumptions about the distribution of the data, and therefore its relevance in real scenarios. In essence, by having data from the past now referred to as training data, we can sort out places based on a particular attribute.

For example, take a look at the following table of data points, which consists of two characteristics:

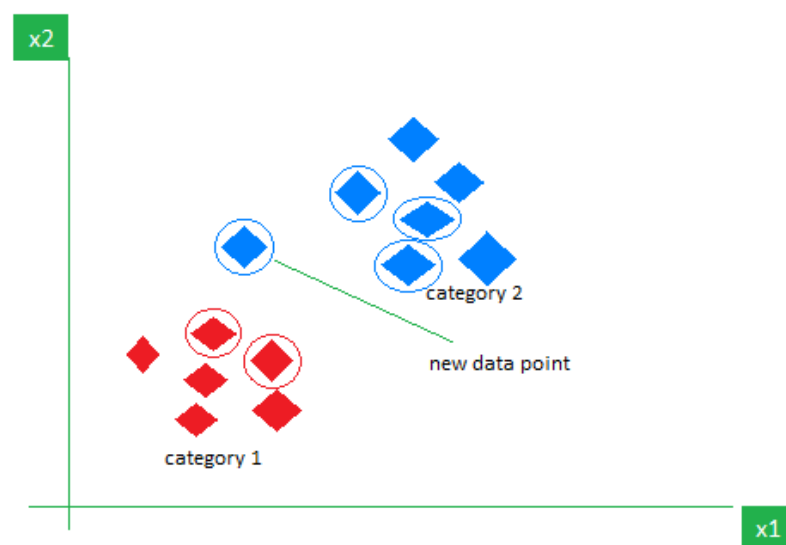


Figure 4. Algorithm working visualization USING KNN [18]

At to this stage, the testing data is given, which also known as a different is set of data points. You will assign these points to a group by considering the training set usually composed of elements similar to the entire computationally intensive process. Kindly note that the ones that do not fall under any or the above groups are labelled “White.”

Distance Metrics Used in KNN Algorithm

Regarding the query location, KNN makes it possible to determine the points, or clusters that are nearest to the query location. However, when aiming to decide which points belong to some cluster and are closest to a query location, a metric is crucial. For the sake of this discussion, we make use of the following distance measurements:

This is the Cartesian distance between two points if two points are placed on a plane or hyperplane. It's the length of the portion of the straight line that joining two points or in case of the distance in between two locations. This means that it is possible for us to determine the net displacement that might take place between two states of an item by use of this measurement. K-NN is a machine learning model among the simplest and easiest to implement, known as the K-nested neural network. In relations of the basic circulation of the information, one does not have to presume anything either. As it can process both quantitative and qualitative data, it can be used when pragmatic to numerous dissimilar kinds of data sets with regard to classification and regression. By comparing the metrics utilized in the database, this novel and simple approach predicts the effectiveness of the study. In this respect, let us note that K-NN is significantly less sensitive to outliers compared with other methods. Similar to the K-NN

approach, the distance metric is determined as the Euclidean distance to a set of K neighbours, using which K-NN decides a number of data points nearest to a given information point. This is done constructed on the identification of the K nearest neighbors. Subsequently, its class or value is computed as the average or the mode of the K nearest neighbors. By so doing, the modulation of the algorithm is made in a way that it is capable of training on many patterns to give prediction depending on the situation of the local data.

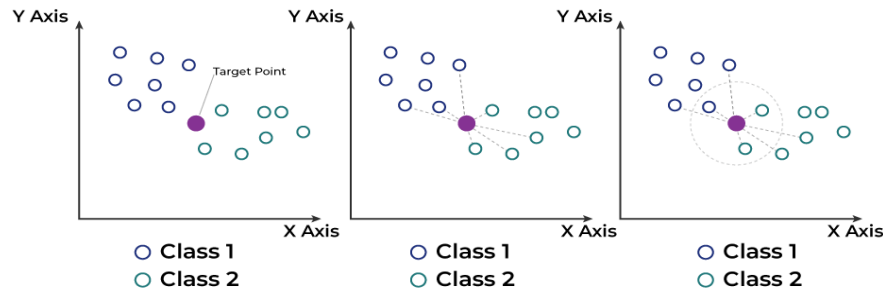


Figure 5. Step process of KNN Algorithm [18]

Step-by-Step clarification of how KNN works is discoursed below:

Step 1: Choosing the appropriate K value

K will be the number of nearest neighbors which one should consider when proceeding to make the prediction.

Step 2: Finding distances

In this case, the distance measurement that is used in order to compare target and training data elements is the Euclidean distance. For the purpose of this computation, the distance is calculated to every other data point from the particular defined target point.

Step 3: Retrieving most similar data points (Nearest Neighbors)

Basic knowledge of K Nearest Neighbors indicates that K Nearest Neighbors to the target point are those which are nearest.

Step 4: Voting for Classification or Averaging for Regression

In situations when there is a problem with categorization, the process of choosing the class level labels among classes is carried out through the utilization of approaches that involve majority voting. For the data point in question, a forecast is formed based on the class that appears the most frequently among the neighbours that have been defined. This class makes up the majority of the neighbours. As part of the process of regression, the class label is determined by computation the regular of the goal values of the K neighbours that are closest to the subject. This is done in order to regulate the class tag. It is the computation of the average value that serves as the outcome that is predicted for the data point that has been provided. In order to denote the training dataset, which will be comprised of n data points, each of which will be defined by a d-dimensional feature vector, the letter X will be utilized. To represent the labels or values that correlate to each data point in X, Y will be utilized as the representation. During the process of computing the detachment amongst each information point in X and the new data point x, the approach makes use of a detachment metric, such as the Euclidean detachment, in order to accurately determine the distance. This step is done after the operation has successfully collected a new data point say x. Once obtained, this procedure helps identify the bases in X, which have K data points closest to x. Within classification concerns, it's applied to estimate which label y is most frequent amongst the closest K neighbours of x. While performing regression tasks the technique works out the average or weighted average of the y's of the K nearest neighbours and then labels that value as the value that is expected to be x. This makes it possible for the method to determine the value of x as projected in completion of the given task.

5. Feature Extraction

The most basic temporal analysis technique in the domain is statistical feature extraction of the data. As the relevance of statistical programming languages continues to grow, this process is improved as a result of the inherent and integrated capabilities that these languages provide. The integration of emotive aptitude and machine intellect into HCI classifications is accomplished by the employment of numerical feature extraction with electroencephalogram gestures. Over the course of many weeks, research was conducted with the objective of recognizing various emotional states through the gathering of EEG data from a large number of participants. This categorization is applicable to a topic that may be clearly specified for the purpose of conducting an objective analysis that is free of any personal prejudice. This is a list of the qualities that were described in the research carried out by Picard et al. in 2001:

1. Mean (raw signal)

$$\mu_X = \frac{1}{N} \sum_{n=1}^N X_n,$$

where X_n denotes the rate of the n th model of the rare signal and $n = 1:N$ information points in the rare indication.

2. Standard deviation (STD) (raw signal)

$$\sigma_X = \left(\frac{1}{N-1} \sum_{n=1}^N (X_n - \mu_X)^2 \right)^{\frac{1}{2}}$$

3. Mean of complete standards of first differences (raw signal)

$$\delta_X = \frac{1}{N-1} \sum_{n=1}^{N-1} |X_{n+1} - X_n|$$

4. Mean of complete standards of first differences (normalized signal)

$$\tilde{\delta}_X = \frac{1}{N-1} \sum_{n=1}^{N-1} |\tilde{X}_{n+1} - \tilde{X}_n| = \frac{\delta_X}{\sigma_X},$$

6. Convolutional Neural Networks

This study assesses the efficiency of CNN methods in computer vision applications relative to machine learning techniques. This is achieved by extracting values from the EEG signal data.

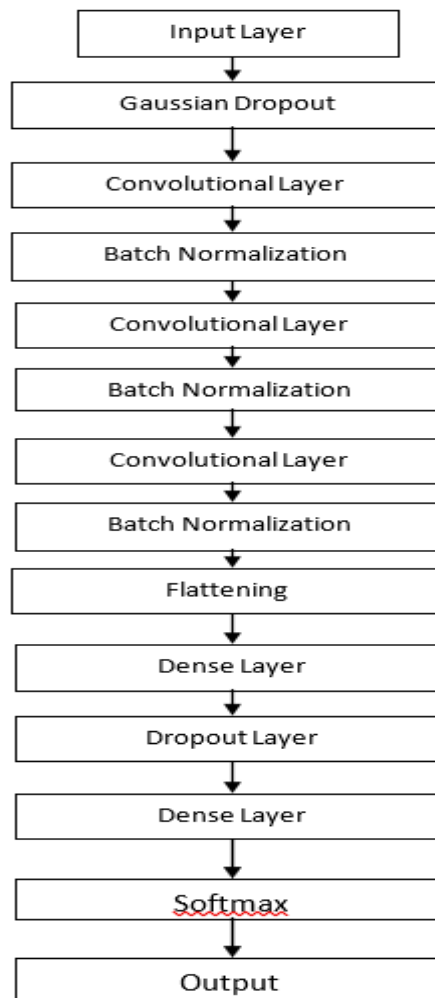


Figure 6. Convolutional neural network

An example of a separate architecture for deep learning is provided by the convolutional neural network model. The implementation of CNNs is quite precise; yet, the recital of the model is substantially impacted by a number of influences, counting the number of layers, dropout rates, and other characteristics. As can be seen in Figure 6, the model is created in this manner. In order to guarantee precision, the training and assessment of DL-CNN necessitate the validation of every picture by means of a series of kernels or filters. The convolutional layers, max pooling, ReLU, SoftMax layer, and fully connected layer are all removed, and the classification layer is used to classify objects that fall within the range of [0,1] based on the probabilistic relevance of their characteristics. The ReLU layer is utilized in networks for the purpose of rectifying the hidden layers. This layer is also known as the ReLU layer. The ReLU function is a basic intention that, if the input rate is greater than zero, returns the value that was previously supplied and, else, returns zero. These are some of the ways in which the function $\max()$ and the variable x can be declared explicitly [19]: y is the maximum of 0 and x (20)

Max pooling layer

Subsampling is utilized by this layer to decrease the number of parameters in bigger photographs. As a result, any function that modifies dimensionality is eliminated while crucial information is preserved. The largest value that can be found inside the adjusted feature map is evaluated using max pooling [19].

PCA (Principal component analysis)

The reduction of dimensionality is accomplished by the application of a method of machine learning known as critical factor analysis. Essential statistical methods and linear algebra matrices are utilized in the process of evaluating source data projections, regardless of whether the projections are being evaluated in a similar or decreased capacity. PCA is a procedure for dimensionality decrease that takes data with ' m ' variables and changes them into a subspace with ' m ' or less dimensions. This approach maintains the most significant aspects of the initial dataset. Let I represent a source image matrix that is n by m , and let J represent the product that is produced. The first thing that has to be done is to figure out the average value of each column. In order to normalize the data in a column, the mean value of the column is subtracted from the values. The final step involves the calculation of the covariance of the centered matrix. It is necessary to provide a list of vector values in order to evaluate the decomposition of each covariance matrix. The orientations or components of J 's reduced subspace are represented by the vectors, while the peak amplitudes of the motions are represented by the other vectors. Utilizing the eigenvalues allows for the evaluation of these vectors, which ultimately results in the creation of a new subspace ranking for the representation. It is common practice to make use of K eigenvectors, which are supposed to represent the essential mechanisms or qualities [19].

Euclidean Distance

A metric is required to compute the detachments between the query word I_y and the retrieved word pictures I_x . Establishing equivalence between query images and names requires a bitwise measurement technique. Thus, the system necessitates a similarity metric in which the detachment value reflects the sum of similar bits in the query imageries [21].

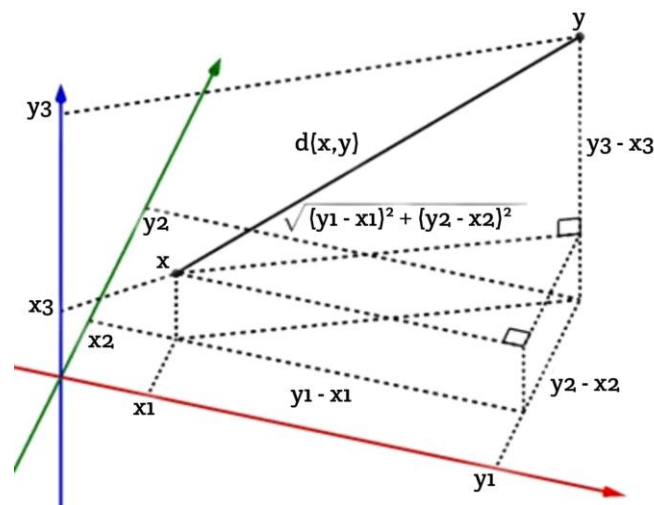


Figure 7. Illustration of Euclidean Distance [21].

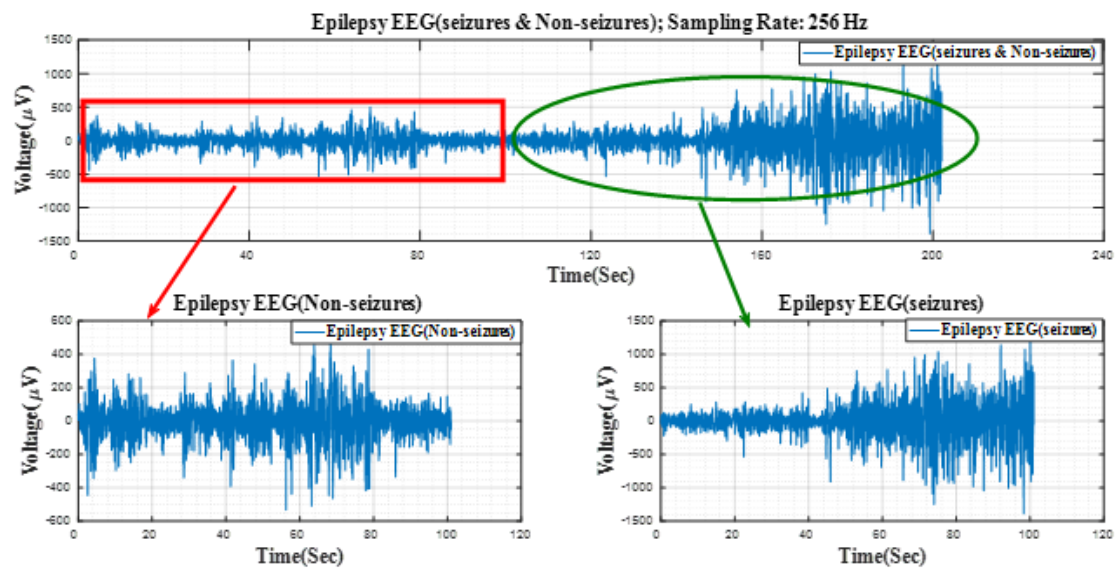


Figure 8. Seizure Detection of EEG Signal

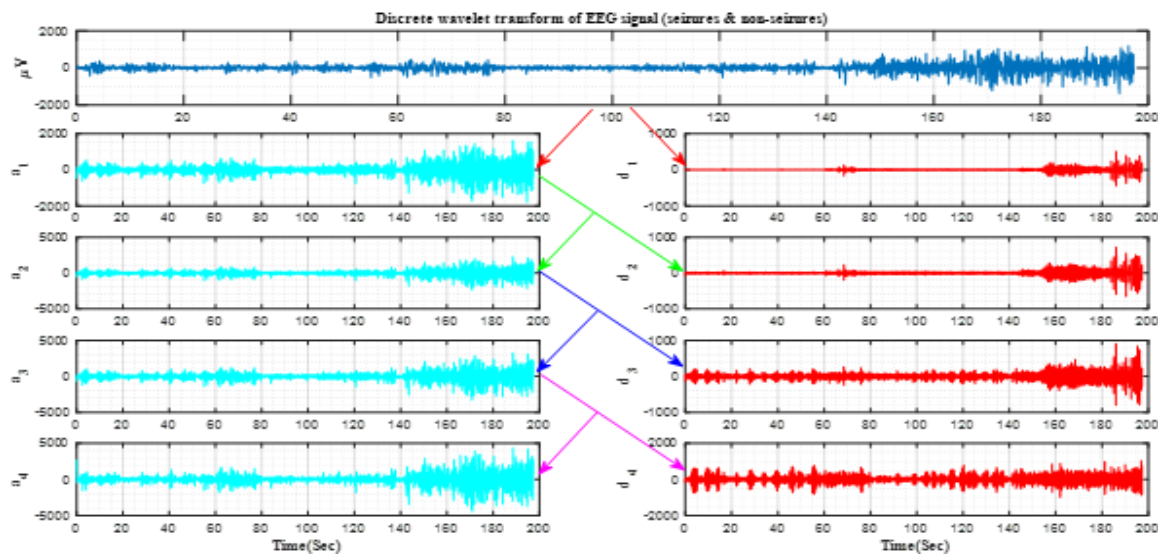


Figure 9. Discrete Wavelet Transform of EEG Signal

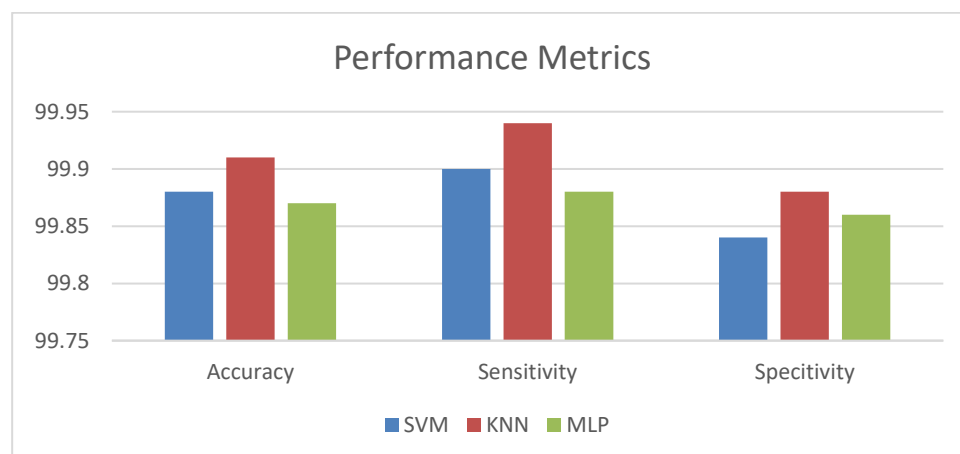


Figure 10. Performance Metrics

Table 1: Performance Metrics

	Accuracy	Sensitivity	Specificity
SVM with Machine Learning	99.88	99.9	99.84
KNN with CNN (Deep learning)	99.91	99.94	99.88
MLP (Machine Learning)	99.87	99.88	99.86

7. Conclusion

When it came to accuracy and latency, the two methods, SVM that utilized machine learning and KNN that utilized DconvNET, both demonstrated outcomes that were comparable to one another. On data that had been standardized to maintain balanced proportions between seizure and non-seizure samples, KNN displayed statistically significant superiority over SVM, reaching an accuracy of 99.88% compared to 99.91% for the approach that was advised. Between the two methods, there was not a discernible difference in latency that could be considered statistically significant. With regard to the identification of epileptic seizures from EEG data, KNN demonstrated a higher level of performance when assessed on normalized data. Nevertheless, in order to further differentiate between the two techniques in reality, a bigger patient cohort is required.

References

- [1] Anthony K Ngugi et al. "Estimation of the burden of active and life-time epilepsy: a meta-analytic approach". In: *Epilepsia* 51.5 (2010), pp. 883–890.
- [2] U. R. Acharya et al. "Automated diagnosis of epileptic EEG using entropies". In: *Biomedical Signal Processing and Control* 7.4 (2012), pp. 401–408.
- [3] A. Jacoby, D. Snape, and G. A. Baker. "Epilepsy and social identity: the stigma of a chronic neurological disorder." In: *The Lancet Neurology* 4.3 (2005), pp. 171–178.
- [4] S. N. Abdulkader, A. Atia, and M. S. M. Mostafa. "Brain computer interfacing: Applications and challenges". In: *Egyptian Informatics Journal* 16.2 (2015), pp. 213–230.
- [5] A. K. Tiwari et al. "Automated diagnosis of epilepsy using key-point based local binary pattern of EEG signals." In: *IEEE journal of biomedical and health informatics* 21.3 (2017), pp. 888–896.
- [6] C. C. Jouney, P. J. Franaszczuk, and G. K. Bergey. "Improving early seizure detection." In: *Epilepsy and behavior* 1 (2011), pp. 44–8.
- [7] Yuedong Song and Pietro Liò. "A new approach for epileptic seizure detection: sample entropy based feature extraction and extreme learning machine". In: *Journal of Biomedical Science and Engineering* 3.06 (2010), p. 556.
- [8] Lachezar Bozhkov et al. "EEG-based subject independent affective computing models". In: *Procedia Computer Science* 53 (2015), pp. 375–382.
- [9] Ali H Shueb and John V Guttag. "Application of machine learning to epileptic seizure detection". In: *Proceedings of the 27th International Conference on Machine Learning (ICML-10)*. 2010, pp. 975–982.
- [10] E. Pippa et al. "Classification of epileptic and non-epileptic EEG events." In: *2014 4th International Conference on Wireless Mobile Communication and Healthcare-Transforming Healthcare Through Innovations in Mobile and Wireless Technologies* (2014), pp. 87–89.
- [11] Ralph Meier et al. "Detecting epileptic seizures in long-term human EEG: a new approach to automatic online and real-time detection and classification of polymorphic seizure patterns". In: *Journal of clinical neurophysiology* 25.3 (2008), pp. 119–131.
- [12] AJ Gabor, RR Leach, and FU Dowla. "Automated seizure detection using a self-organizing neural network". In: *Electroencephalography and clinical Neurophysiology* 99.3 (1996), pp. 257–266.
- [13] Adam Page et al. "A flexible multichannel EEG feature extractor and classifier for seizure detection". In: *IEEE Transactions on Circuits and Systems II: Express Briefs* 62.2 (2015), pp. 109–113.
- [14] Scott B Wilson et al. "Seizure detection: evaluation of the Reveal algorithm". In: *Clinical neurophysiology* 115.10 (2004), pp. 2280–2291.
- [15] A. L. Goldberger et al. "PhysioBank, PhysioToolkit, and PhysioNet: Components of a New Research Resource for Complex Physiologic Signals". In: *Circulation* 101.23 (2000 (June 13)), e215–e220.
- [16] Robert S Fisher et al. "Epileptic seizures and epilepsy: definitions proposed by the International League Against Epilepsy (ILAE) and the International Bureau for Epilepsy (IBE)". In: *Epilepsia* 46.4 (2005), pp. 470–472.
- [17] Turkey Alotaiby et al. "A review of channel selection algorithms for EEG signal processing". In: *EURASIP Journal on Advances in Signal Processing* 2015.1 (2015), p. 66.

- [18] <https://www.geeksforgeeks.org/k-nearest-neighbours/>
- [19] S. Albawi, T. A. Mohammed and S. Al-Zawi, "Understanding of a convolutional neural network," 2017 International Conference on Engineering and Technology (ICET), 2017, pp. 1-6, doi: 10.1109/ICEngTechnol.2017.8308186.
- [20] M. Vidya and M. V. Karki, "Skin Cancer Detection using Machine Learning Techniques," 2020 IEEE International Conference on Electronics, Computing and Communication Technologies (CONECCT), 2020, pp. 1-5, doi: 10.1109/CONECCT50063.2020.9198489.
- [21] https://en.wikipedia.org/wiki/Euclidean_distance