

**SONA COLLEGE OF
TECHNOLOGY**
Learning is a Celebration!

| An Autonomous Institution |

**Department of Electronics and
Communication Engineering**



TRANSFORMS AND ALGORITHM IN SIGNAL AND IMAGE PROCESSING

2022-23



Junction Main Road
Suramangalam (PO)
Salem-636 005. TN. India
www.sonatech.ac.in

Editorial Head

Dr.R.S.Sabeenian,
Professor &Head, Dept of ECE,
Head R&D Sona SIPRO



Editorial Members

Dr.M.E.Paramasivam
Associate
Professor



Dr. T.Shanthi
Assistant
(Sr.G)
Professor



Dr.P.M.Dinesh
Assistant
Professor



Prof.R.Anand
Assistant
Professor



Prof.Eldho paul
Assistant Professor



Magazine
co-ordinator
Dr.K.Manju
Assistant
Professor



PREFACE

The field of signal and image processing encompasses the theory and practice of algorithms and hardware that convert signals produced by artificial or natural means into a form useful for a specific purpose. The signals might be speech, audio, images, video, sensor data, telemetry, electrocardiograms, or seismic data, among others; possible purposes include transmission, display, storage, interpretation, classification, segmentation, or diagnosis.

Current research in digital signal processing includes robust and low complexity filter design, signal reconstruction, filter bank theory, and wavelets. In statistical signal processing, the areas of research include adaptive filtering, learning algorithms for neural networks, spectrum estimation and modeling, and sensor array processing with applications in sonar and radar. Image processing work is in restoration, compression, quality evaluation, computer vision, and medical imaging. Speech processing research includes modeling, compression, and recognition. Video compression, analysis, and processing projects include error concealment technique for 3D compressed video, automated and distributed crowd analytics, stereo-to-auto stereoscopic 3D video conversion, virtual and augmented reality.

BRAIN TUMOR DETECTION AND ENHANCEMENT USING DEEP LEARNING

KISHAN RAJ S, PAVITHRAN RAJ Z,SNEHA S

ABSTRACT

In this project we deal with the implementation of simple algorithm for detection of range and shape of tumor in brain MR images. Tumor is an uncontrolled growth of tissues in any part of the body. Tumors are of different types and they have different characteristics and different treatment. As it is known, brain tumor is inherently serious and life-threatening because of its character in the limited space of the intracranial cavity. Most Research in developed countries show that the number of people who have brain tumors were died due to the fact of inaccurate detection. Generally, CT scan or MRI that is directed into intracranial cavity produces a complete image of brain. This image is visually examined by the physician for detection & diagnosis of brain tumor. However, this method of detection resists the accurate determination of stage & size of tumor. To avoid that, this project uses computer aided method for detection of brain tumor based on the combination of two algorithms. This method allows the segmentation of tumor tissue with accuracy and reproducibility comparable to manual segmentation. In addition, it also reduces the time for analysis. At the end of the process the tumor is extracted from the MR image and its exact position and the shape also determined. The stage of the tumor is displayed based on the amount of area calculated from the cluster.

1.INTRODUCTION

The key issue was detection of brain tumor in very early stages so that proper treatment can be adopted. Based on this information, the most suitable therapy, radiation, surgery or chemotherapy can be decided. As a result, it is evident that the chances of survival of a tumor-infected patient can be increased significantly if the tumor is detected accurately in its early stage. The segmentation was employed to determine the affected tumor part using imaging modalities. Segmentation is process of dividing the image to its constituent parts sharing identical properties such as color, texture, contrast and boundaries. According to world health organization, the grading system scales are used from grade I to grade IV. These grades classify benign and malignant tumor types. The grade I and II are low-level grade tumors while grade III and IV are high-level grade tumors. Brain tumor can affect individuals at any age. The impact on every individual may not be same. Due to such a complex structure of human brain, a diagnosis of tumor area in brain is challenging task. The malignant-type grade III and IV of tumor is fast growing. Affects the healthy brain cells and may spread to other parts of the brain

or spinal cord and is more harmful and may remain untreated. So detection of such brain tumor location, identification and classification in earlier stage is a serious issue in medical science. The brain tumor is an abnormal growth of uncontrolled cancerous tissues in the brain. A brain tumor can be benign and malignant. The benign tumor has uniformity structures and contains non-active cancer cells. The malignant tumor has non-uniformity structures and contains active cancer cells that spread all over parts. By enhancing the new imaging techniques, it helps the doctors to observe and track the occurrence and growth of tumor-affected regions at different stages so that they can take provide suitable diagnosis with these images scanning.

MR Brain Image Segmentation:

The primarily preferred artificial neural networks (ANNs) for medical imaging applications. However, there are some hidden drawbacks associated with conventional LVQ which often go unnoticed. One of the significant drawbacks is the lack of convergence condition which forces the LVQ to completely depend on iterations. Any iteration dependent ANN becomes less accurate since the correct fixation of the number of iterations is extremely difficult. If the number of iterations is not optimal, then the LVQ may encounter local minima problems. In this work, this specific problem is tackled by proposing a hybrid swarm intelligence-LVQ approach in the context of MR brain image segmentation. The PSO is used to train the LVQ which eliminates the iteration-dependent nature of LVQ. The proposed methodology is used to detect the tumor regions in the abnormal MR brain images.

Problem Statement:

The feasibility of these algorithms for analyzing is presented through experimental investigation. The simulation results give that the proposed optimal approach Conclusion Recommendation. The performance of the proposed study is compared with the existing traditional algorithm and real time medical diagnosis image. The diffusion phase, each inactive agent randomly selects another agent from the population; if the selected agent is active, the selecting agent adopts the hypothesis of the active agent and the information sharing takes place.

1. Diffusion radius
2. Number of iterations

2.METHODOLOGY

Proposed System

The MRI or CT scan images are primary follow up diagnostic tools when a neurologic exam indicates a possibility of a primary or metastatic brain tumor existence. The tumor tissue mainly appears in brighter colors than the rest of the regions in the brain. Based on this observation, an automated algorithm for brain tumor detection and medical doctors' assistance in facilitated and accelerated diagnosis procedure has been developed and initially tested on images obtained from the patients with diagnosed tumors and healthy subjects.

Scope of the Project

The consists of some problem in selecting the parameter configuration. In our future study, we will investigate better and more efficient ways to solve the computational problems. Our goal is to achieve real time interactive image segmentation of arbitrary number of classes using the optimization frame work with less computational time.

1. Preprocessing:

The preprocessing step improves the standard of the brain tumor MR images and makes these images suited for future processing by clinical experts or imaging modalities. It also helps in improving parameters of MR images. The parameters includes improvement in signal-to-noise ratio, enhancement in visual appearance of MR images, the removal of irrelevant noise and background of undesired parts, smoothing regions of inner part, maintaining relevant edges.

2. Segmentation:

The segmentation is a process where the image is partitioned into different regions. Let an entire region of image be represented by S . Segmentation process can be viewed as partition of S into p subregions like $S_1, S_2, S_3, \dots, S_p$. Certain conditions has to satisfied such as the segmentation must be intact; that is each and every pixel should be within the region, every points in the regions should be connected in some sense, regions should be disjoint, etc.

3. Thresholding method:

The thresholding method is a basic and powerful method to segment the task in low-contrast images. Histogram analysis is used to select threshold required objects and the selection of an optimized threshold is a difficult values based on image intensity. Thresholding methods are classified into local and global. If high homogeneous contrast

segmentation and many distinct regions are segmented within the gray-by Gaussian distribution method. These methods are utilized when the threshold value cannot be measured from the whole image histogram or single value of the threshold does not provide good results of segmentation. or intensity exists among the objects and background, then the global thresholding method is the best In most cases, the thresholding method is applied at the first stage for level images.

4. Region Growing:

Region growing is grouping of pixels or subregions into larger regions based on certain criteria. The main aim was to select a ‘seed’ points and attach each of these seed to those neighboring pixels having identical properties to grow region. A set of seeds was taken as input within the image and marked the objects to be segmented. The region grows iteratively by estimating all unallocated neighboring pixels of the region. The similarity was the measure of difference between pixel’s intensity value and the region’s mean, δ . The pixel with the smallest difference measured this way was allocated to the respective region. This was continued until all pixels were allocated to a region. Seeded region growing requires seeds as additional input. The results depend on the selection of seeds [13]. The measurement was based on mean value of the pixel intensity. The image gets segmented; this image was used to identify the desired tumor region.

5. Morphological Operations:

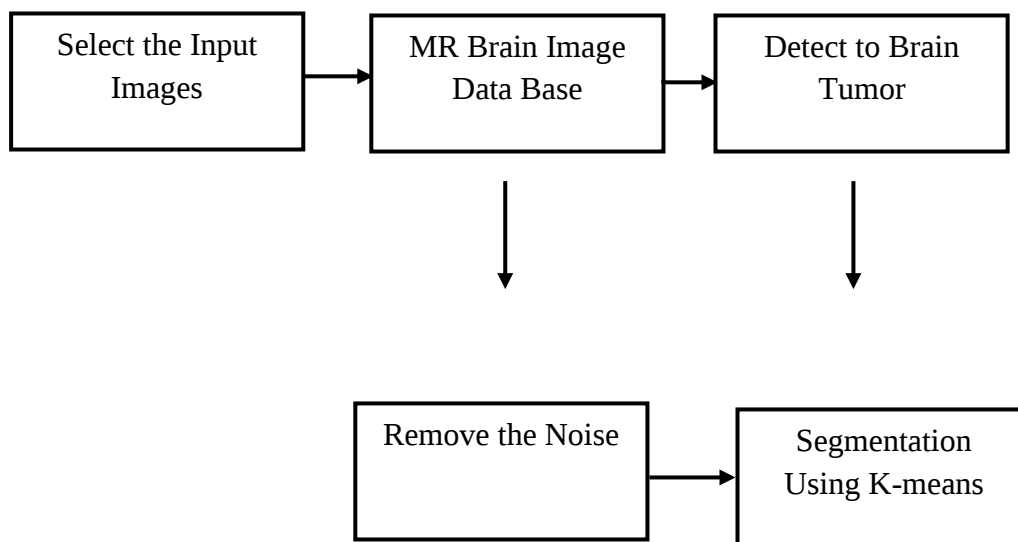
Morphology deals with study of shapes and boundary area extraction from brain tumor images. Morphological operation is rearranging the order of pixel values. It operates on structuring element and input images. Structuring elements are attributes that probes a features of interest. The basic operations used here are dilation and erosion. Dilation operation adds the pixels to boundary region, while erosion removes the pixels from the boundary region of the objects. These operations were carried out based on the structuring elements. Dilation chooses highest value by comparing all pixel values in neighborhood of input image described by structuring element, whereas erosion chooses the lowest value by comparing all the pixel values in the neighborhood of the input image.

Existing System

There are presenting a unique approach by deploying SDS to use in assessing medical

images. This approach demonstrates a promising ability to undertake this task with a similar level of sensitivity. Each scan used in this paper is processed by the SDS agents, which are responsible for locating the desired areas. The reproducibility and the accuracy of the SDS algorithm can be utilised in developing a standardised system to help interpreting medical images and prevent operator errors and discrepancies. This type of technologies can be employed as an adjunct to help radiologists assess the various types of images making the diagnosis more thorough and less time consuming. this work highlights the application of swarm intelligence techniques for MR brain segmentation.

Existing Block Diagram



Block diagram of Brain Tumor Detection

In this Block diagram, the selected input are collected in MRI format ,first the selected input images are send through the MR Brain Image Data Base, here we have to do two process one is removing the noise of the image and to detect the brain tumor using Deep learning .If the image reduces the noise there will be good quality of the image, and parallel the detection is to be done, after the both processes then it takes to segmentation process to detect the brain tumor and points out the defect where it is having the tumor. These days issue of brain tumor automatic identification is of great interest. Tumor is unusual growth of the tissues. A brain tumor is a quantity of unnecessary cells growing in the brain or central spine canal. Here we use the best CT scan.

Implementation and Result Discussion

This study proposes an umbrella deployment of swarm intelligence algorithm, such as

stochastic diffusion search for medical imaging applications. After summarising the results of some previous works which shows how the algorithm assists in the identification of metastasis in bone scans and microcalcifications on mammographs, for the first time, the use of the algorithm in assessing the CT images of the aorta is demonstrated along with its performance in detecting the nasogastric tube in chest X-ray. The swarm intelligence algorithm presented in this study is adapted to address these particular tasks and its functionality is investigated by running the swarms on sample CT images and X-rays whose status have been determined by senior radiologists. In addition, a hybrid swarm intelligence-learning vector quantisation (LVQ) approach is proposed in the context of magnetic resonance (MR) brain image segmentation. The particle swarm optimisation is used to train the LVQ which eliminates the iteration-dependent nature of LVQ. The proposed methodology is used to detect the tumour regions in the abnormal MR brain images.

The term image registration basically denotes the process of alignment of images. In recent years, the rapid growth in the field of military automatic target recognition remote cartography, computer vision, image fusion, medical imaging, and astrophotography has established the need for the development of good image registration technique. Image registration is a process in which final information is gained from different data sources.

It is a process of aligning two images acquired by same/different sensors, at different times or from different viewpoint. This paper presents the image registration techniques based on extracting interest point area of scene images using Discrete wavelet Transform and PSO. The root mean square error is used as similarity measure for finding out the area of interest. The usual method to detect brain tumor is Magnetic Resonance Imaging (MRI) scans. From the MRI images information about the abnormal tissue growth in the brain is identified. In various research papers, the detection of brain tumor is done by applying Machine Learning and Deep Learning algorithms. The purpose of this paper is to investigate the use of biologically-inspired swarm methods for signal filtering. The signal, in the case of images the grayscale value of the pixels along a line in the image, is modeled by the trajectory of an agent playing the role of the prey for a swarm of hunting agents. The swarm hunting the prey is the system performing the signal processing. Several models that try to find accurate and efficient boundary curves of brain tumors in medical images have been implemented in the literature.

3.PROJECT DESCRIPTION

Input Design

The input design is the link between the information system and the user. It comprises the developing specification and procedures for data preparation and those steps are necessary to put transaction data in to a usable form for processing can be achieved by inspecting the computer to read data from a written or printed document or it can occur by having people keying the data directly into the system.

The design of input focuses on controlling the amount of input required, controlling the errors, avoiding delay, avoiding extra steps and keeping the process simple. The input is designed in such a way so that it provides security and ease of use with retaining the privacy. Input Design considered the following things:

1. What data should be given as input?
2. How the data should be arranged or coded?
3. The dialog to guide the operating personnel in providing input.
4. Methods for preparing input validations and steps to follow when error occur.

Objectives

1. Input Design is the process of converting a user-oriented description of the input into a computer-based system. This design is important to avoid errors in the data input process and show the correct direction to the management for getting correct information from the computerized system.
2. It is achieved by creating user-friendly screens for the data entry to handle large volume of data. The goal of designing input is to make data entry easier and to be free from errors. The data entry screen is designed in such a way that all the data manipulates can be performed. It also provides record viewing facilities.
3. When the data is entered it will check for its validity. Data can be entered with the help of screens. Appropriate messages are provided as when needed so that the user will not be in maize of instant. Thus the objective of input design is to create an input layout that is easy to follow

Output Design

A quality output is one, which meets the requirements of the end user and presents the information clearly. In any system results of processing are communicated to the users and to other system through outputs. In output design it is determined how the information is to be displaced for immediate need and also the hard copy output. It is the most important and direct source information to the user. Efficient and intelligent output design improves the system's relationship to help user decision-making.

1. Designing computer output should proceed in an organized, well thought out manner; the right output must be developed while ensuring that each output element is designed so that people will find the system can use easily.
2. Select methods for presenting information.
3. Create document, report, or other formats that contain information produced by the system.

The output form of an information system should accomplish one or more of the following objectives.

1. Convey information about past activities, current status or projections of the Future.
2. Signal important events, opportunities, problems, or warnings.
3. Trigger an action.
4. Confirm an action.

Feasibility Study

The feasibility of the project is analyzed in this phase and business proposal is put forth with a very general plan for the project and some cost estimates. During system analysis the feasibility study of the proposed system is to be carried out.

This is to ensure that the proposed system is not a burden to the company. For feasibility analysis, some understanding of the major requirements for the system is essential.

Three key considerations involved in the feasibility analysis are

5. Economical Feasibility.
6. Technical Feasibility.
7. Social Feasibility.

Economical Feasibility

This study is carried out to check the economic impact that the system will have on the organization. The amount of fund that the company can pour into the research and development of the system is limited. The expenditures must be justified. Thus the developed system as well within the budget and this was achieved because most of the technologies used are freely available. Only the customized products had to be purchased.

Technical Feasibility

This study is carried out to check the technical feasibility, that is, the technical requirements of the system. Any system developed must not have a high demand on the available technical resources. This will lead to high demands on the available technical resources. This will lead to high demands being placed on the client. The developed system must have a modest requirement, as only minimal or null changes are required for implementing this system.

Social Feasibility

The aspect of study is to check the level of acceptance of the system by the user. This includes the process of training the user to use the system efficiently. The user must not feel threatened by the system, instead must accept it as a necessity. The level of acceptance by the users solely depends on the methods that are employed to educate the user about the system and to make him familiar with it.

Future Enhancement

As a future work, the proposed approach will be compared with modified approaches such as adaptive chaotic PSO and hybrid kernel support vector machine. Metastasis to the brain is the most feared complication of systemic cancer and the most common intracranial tumor in adults. Different classifiers can be used to increase the accuracy combining more efficient segmentation and feature extraction techniques with real- and clinical-based cases by using large dataset covering different scenarios.

In future, this technique can be developed to classify the tumors based on Extraction. This technique can be applied to various parts of human body to detect The tumor. Different types of filters and algorithms can be used for further increase in accuracy of the tumor segmentation and detection.

4.RESULT AND DISCUSSIONS

Histogram Equalization

Histogram Equalization is a computer image processing technique used to improve contrast in images. It accomplishes this by effectively spreading out the most frequent intensity values, i.e. stretching out the intensity range of the image. This method usually increases the global contrast of images when its usable data is represented by close contrast values. This allows for areas of lower local contrast to gain a higher contrast.

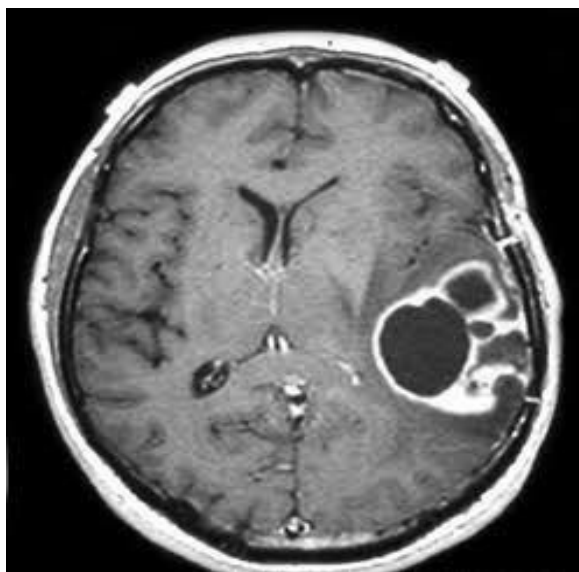


Fig.6.1.a Input image

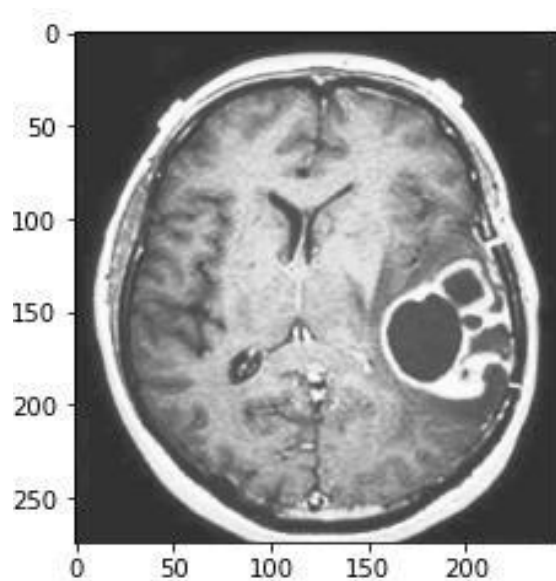


Fig.6.1.b Histogram Equalized Output
image

Histogram equalization is used to enhance contrast. It is not necessary that contrast will always be increase in this. There may be some cases were histogram equalization can be worse. In that cases the contrast is decreased.

Image negative

Image is also known as a set of pixels. When we store an image in computers or digitally, its corresponding pixel values are stored. So, when we read an image to a variable using OpenCV in Python, the variable stores the pixel values of the image. When we try to negatively transform an image, the brightest areas are transformed into the darkest and the darkest areas are transformed into the brightest.



Fig.6.2.a Input image

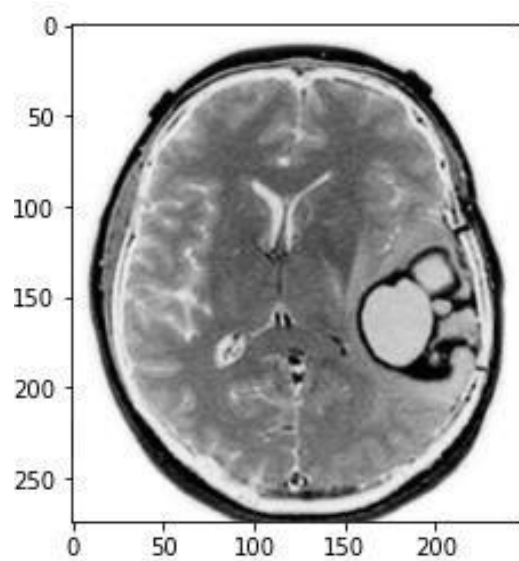


Fig.6.2.b Image negative

Output image

Inverting a digital image is a point processing operation. The output of image inversion is a negative of a digital image. Image inversion or Image negative helps finding the details from the darker regions of the image. Image negative is produced by subtracting each pixel from the maximum intensity value.

Logarithmic Transformation

Logarithmic transformation of an image is one of the gray level image transformations. Log transformation of an image means replacing all pixel values, present in the image, with its logarithmic values. Log transformation is used for image enhancement as it expands dark pixels of the image as compared to higher pixel values.

Log transformation of gives actual information by enhancing the image. If we apply this method in an image having higher pixel values then it will enhance the image more and actual information of the image will be lost. So, this method can't be applied everywhere. It can be applied in images where low pixel values are more than higher ones.

Piecewise Linear transformation functions

Piece-wise Linear Transformation is type of gray level transformation that is used for image enhancement. It is a spatial domain method. It is used for manipulation of an image so that the result is more suitable than the original for a specific application.



Fig.6.3.a Input image

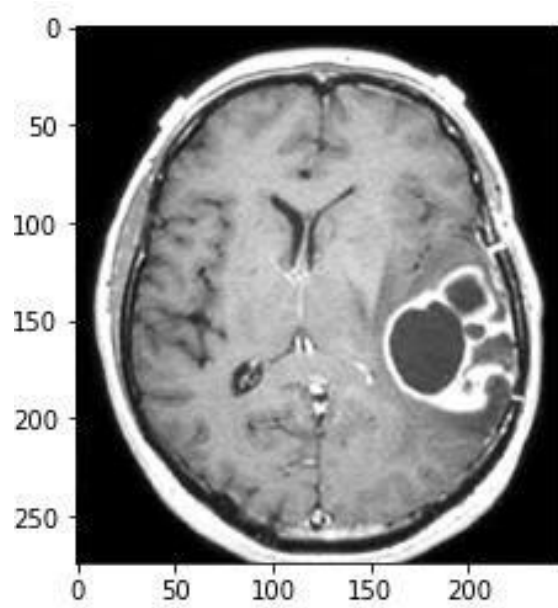


Fig.6.3.b Logarithmic Transformed
Output image

One of the most commonly used piecewise-linear transformation functions is contrast stretching. This process expands the range of intensity levels in an image so that it spans the full intensity of the camera/display.

Power Law transformation

LOG transform enhances small magnitude input values into wider range of output pixel values and compresses large magnitude input values into narrow range of output values. It is useful to display Fourier Transformed images, The drawback for this transform is that the transformation function is fixed and cannot be changed as per requirement

It shows the enhanced images for different type of pixel values. Gamma correction is important for displaying images on a screen correctly, to prevent bleaching or darkening of images when viewed from different types of monitors with different display settings. This is done because our eyes perceive images in a gamma-shaped curve, whereas cameras capture images in a linear fashion.

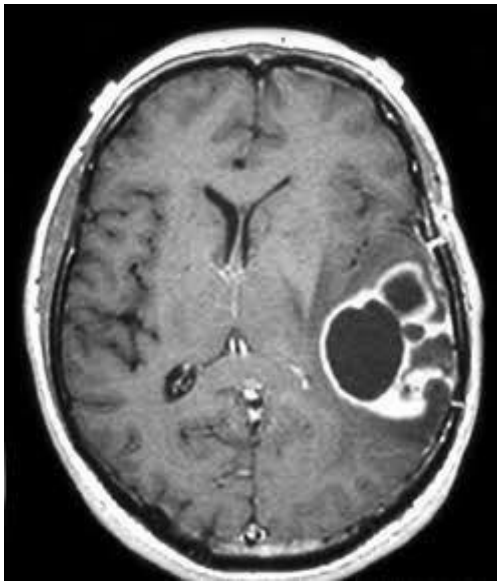


Fig.6.4.a Input image

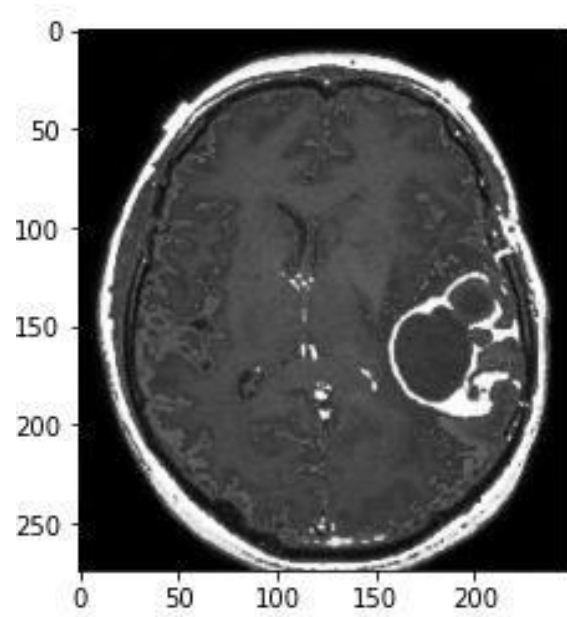


Fig.6.4.b Piecewise linear transformed
Output image

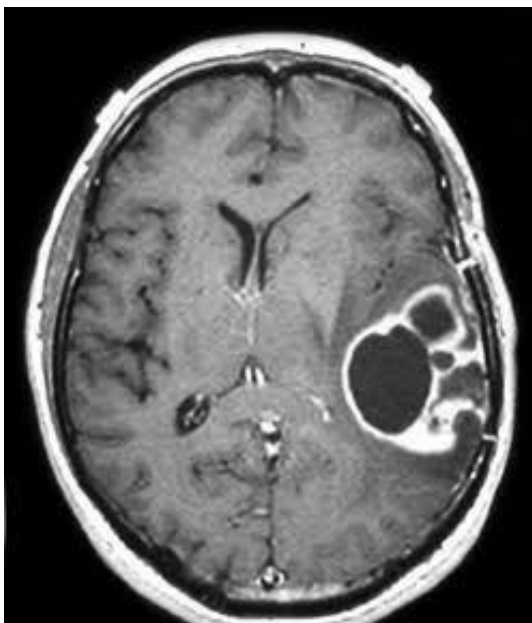


Fig.6.5.a Input image

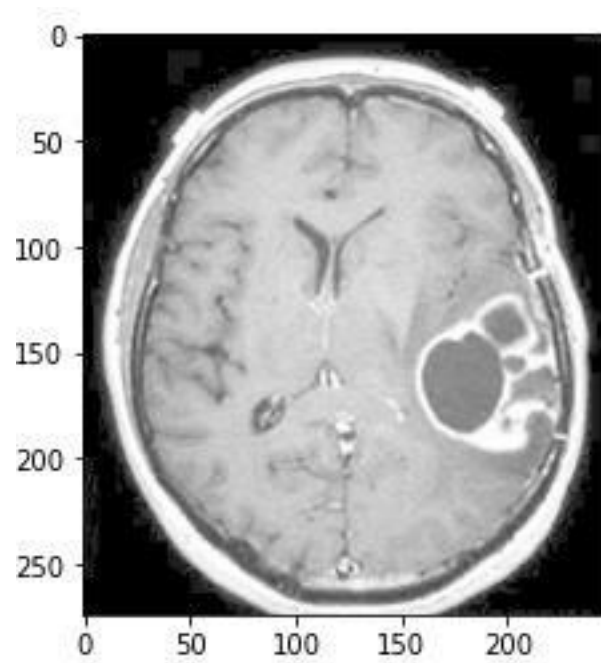


Fig.6.5.b Power law Transformed
Output image

Thresholding transformation

Thresholding is a technique in OpenCV, which is the assignment of pixel values in relation to the threshold value provided. In thresholding, each pixel value is compared with the threshold value. If the pixel value is smaller than the threshold, it is set to 0, otherwise, it is set to a maximum value. Thresholding is a very popular segmentation technique, used for separating an object considered as a foreground from its background. A threshold is a value which has two regions on its either side. below the threshold or above the threshold.



Fig.6.6.a Input image

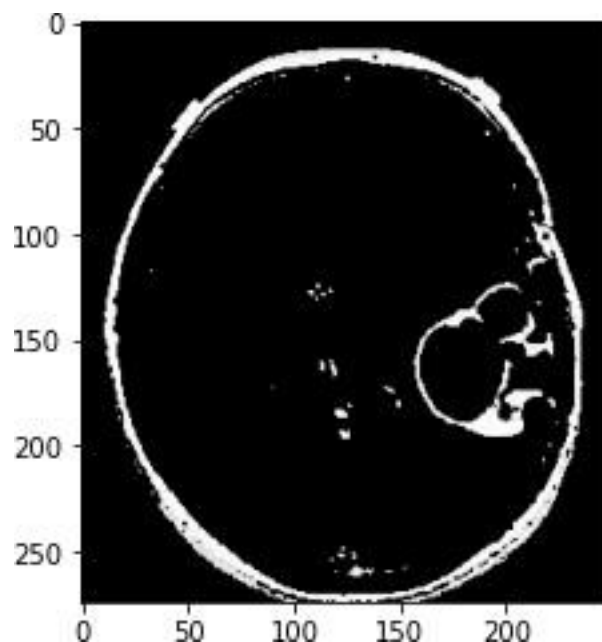


Fig.6.6.b Thresholding transformation

Output image

Here, the brain tumor image is detected it shown in the output image it points out the defect of tumor where it is present.

CONCLUSION

In this research, we have used brain MR images, segmented into normal brain tissue (unaffected) and abnormal tumor tissue (infected). To remove a noise and smoothen the image, preprocessing is used which also results in the improvement of signal-to-noise ratio. Next, we have used discrete wavelet transform that decomposes the images and textural features were

extracted from gray-level co- occurrence matrix (GLCM) followed by morphological operation. Probabilistic neural network (PNN) classifier is used for the classification of tumors from brain MRI images. From the observation results, it can be clearly expressed that the detection of brain tumor is fast and accurate when compared to the manual detection carried out by clinical experts. The performance factors evaluated also shows that it gives better outcome by improving PSNR and MSE parameters. The proposed methodology results in accurate and speedy detection of tumor in brain along with identification of precise location of the tumor. In identification and classification into normal and abnormal tumors from brain MR images, accuracy of nearly 100% was achieved for trained dataset because the statistical textural features were extracted from LL and HL subbands wavelet decomposition and 95% was achieved for tested dataset. With the above results, we conclude that our proposed method clearly distinguishes the tumor into normal and abnormal which helps in taking clear diagnosis decisions by clinical experts.

EEG SIGNAL CLASSIFICATION USING CONVOLUTIONAL NEURAL NETWORK

NISHANTH S, RAJKUMAR S

ABSTRACT

One of the challenges in modeling cognitive events from electroencephalogram (EEG) data is finding representations that are invariant to inter- and intra-subject differences, as well as to inherent noise associated with such data. Herein, we propose a novel approach for learning such representations from multi-channel EEG time-series and demonstrate its advantages in the context of mental load classification task. First, we transform EEG activities into a sequence of topology preserving multi-spectral images, as opposed to standard EEG analysis techniques that ignore such spatial information. Next, we train a deep recurrent-convolutional network inspired by state-of-the-art video classification to learn robust representations from the sequence of images. The proposed approach is designed to preserve the spatial, spectral, and temporal structure of EEG which leads to finding features that are less sensitive to variations and distortions within each dimension. Empirical evaluation on the cognitive load classification task demonstrated significant improvements in classification accuracy over current state-of-the-art approaches in this field.

1. INTRODUCTION

Deep neural networks have recently achieved great success in recognition tasks within a wide range of applications including images, videos, speech, and text (Krizhevsky et al., 2012; Graves et al., 2013; Karpathy & Toderici, 2014; Zhang & LeCun, 2015; Hermann et al., 2015). Convolutional neural networks (ConvNets) lie at the core of best current architectures working with images and video data, primarily due to their ability to extract representations that are robust to partial translation and deformation of input patterns (LeCun et al., 1998). On the other hand, recurrent neural networks have delivered state-of-the-art performance in many applications involving dynamics in temporal sequences [1], such as, for example, handwriting and speech recognition (Graves et al., 2013; 2008). In addition, combination of these two network types have recently been used for video classification (Ng & Hausknecht, 2015). Despite numerous successful applications of deep neural networks to large-scale image, video and text data, they remain relatively unexplored in neuroimaging domain. Perhaps one of the main reasons here is that the number of samples in most neuroimaging datasets is limited, thus making such data less adequate for training large-scale networks with millions of parameters. As it is often demonstrated, the advantages of deep neural networks over traditional machine-learning techniques become more apparent when the dataset size becomes very large. Nevertheless, deep belief network and

Convnets have been used to learn representations from functional Magnetic Resonance Imaging (fMRI) and Electroencephalogram (EEG) in some previous work with moderate dataset sizes (Plis et al., 2014; Mirowski et al., 2009). Plis et al. (2014) showed that adding several Restricted Boltzmann Machine layers to a deep belief network and using supervised pretraining results in networks that can learn increasingly complex representation of the data and achieve considerable increase in classification accuracy. In other works, convolutional and recurrent neural networks have been used to extract representations from EEG time series (Cecotti & Graser, 2011; Guler et al., 2005). For instance, Convnets were used to encode the 2-D structure of EEG feature maps (Mirowski et al., 2009). These studies demonstrated potential benefits of adopting (down-scaled) deep neural networks in neuroimaging, even in the absence of extremely large, million-sample datasets, such as those available for images, video, and text modalities [2]. Herein, we explore the capabilities of deep neural nets for modeling cognitive events from EEG data. EEG is a widely used noninvasive neuroimaging modality which operates by measuring changes in electrical voltage on the scalp induced from cortical activity. Using the classical blind-source separation analogy, EEG data can be thought of as a multi-channel “speech” signal obtained from several “microphones” (associated with EEG electrodes) that record signals from multiple “speakers” (that correspond to activity in cortical regions). State-of-the-art mental state recognition using EEG consists of manual feature selection from continuous time series and applying supervised learning algorithms to learn the discriminative manifold between the states (Lotte & Congedo, 2007; Subasi & Ismail Gursay, 2010).

A key challenge in correctly recognizing mental states from observed brain activity [3] is constructing a model that is robust to translation and deformation of signal in space, frequency, and time, due to inter- and intra-subject differences, as well as signal acquisition protocols. Much of the variations originate from slight individual differences in cortical mapping and/or functioning, giving rise to observed differences in spatial, spectral, and temporal patterns. Moreover, EEG caps which are used to place the electrodes on top of predetermined cortical regions can be another source of spatial variations in observed responses due to imperfect fitting of the cap on heads of different sizes and shapes. An example illustrating potentially high inter- and intra-subject variability in EEG data is given in Appendix. We propose a novel approach to learning representations [4] from EEG data that relies on deep learning and appears to be more robust to inter- and intra-subject differences, as well as to measurement related noise. Our approach is fundamentally different from the previous attempts to learn high level representations from EEG using deep neural networks. Specifically, rather than representing low-level EEG features as a vector, we transform the data into a multi-dimensional tensor which retains the structure of the data throughout the learning process. In other words, we obtain a sequence of topology-preserving multi-spectral images, as opposed to standard EEG analysis

techniques that ignore such spatial information.

Objective of Research:

The primary goal was to develop a classifier that can correctly identify whether a subject is visualizing a task that is familiar or unfamiliar. Secondary goals included providing insight into which brain regions and frequency bands associate with each of the respective classes. If a deep learning approach is found to be viable, these insights may correspond to latent features found within the neural network. Other insights may be obtained from more traditional data processing and machine learning techniques.

2.PROBLEM STATEMENT AND DATA COLLECTION

An ElectroEncephaloGram (EEG) is a test that detects electrical activity in your brain using small, flat metal discs (electrodes) attached to your scalp. Your brain cells communicate via electrical impulses and are active all the time, even when you're asleep. This activity shows up as wavy lines on an EEG recording. [Mayo Clinic. The goal of this project was to classify brain states from EEG data [28]. A joint CU Anschutz/ ULN project has collected EEG data on subjects during sessions in which the subjects were instructed to visualize performing a motor- based task.

The primary goal was to develop a classifier that can correctly identify whether a subject is visualizing a task that is familiar or unfamiliar. Secondary goals included providing insight into which brain regions and frequency bands associate with each of the respective classes. If a deep learning approach is found to be viable, these insights may correspond to latent features found within the neural network. Other insights may be obtained from more traditional data processing and machine learning techniques.

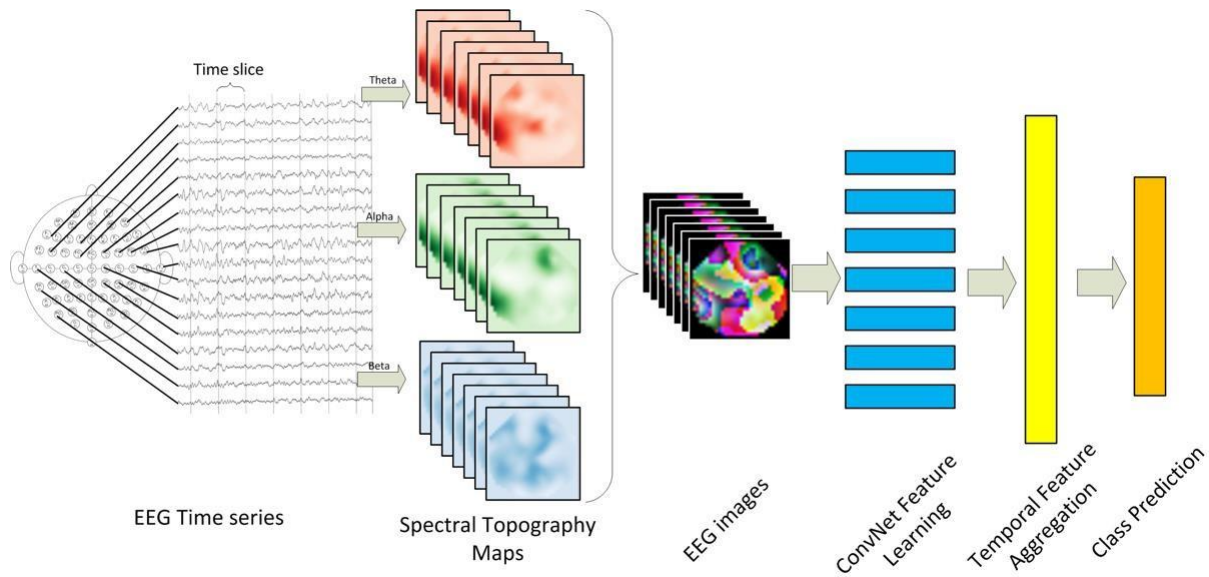
3. METHODOLOGY

Relying on previous EEG research done by Beshivan et. al. [1], as well as the latest advances in video classification [3], the approach was to process the 14- channel time-series data into discrete one-second 'frames' and project these frames onto a 2D map of the surface of the head. Then a convolutional neural network (CNN) was trained to classify frames as shown in figure shown below.

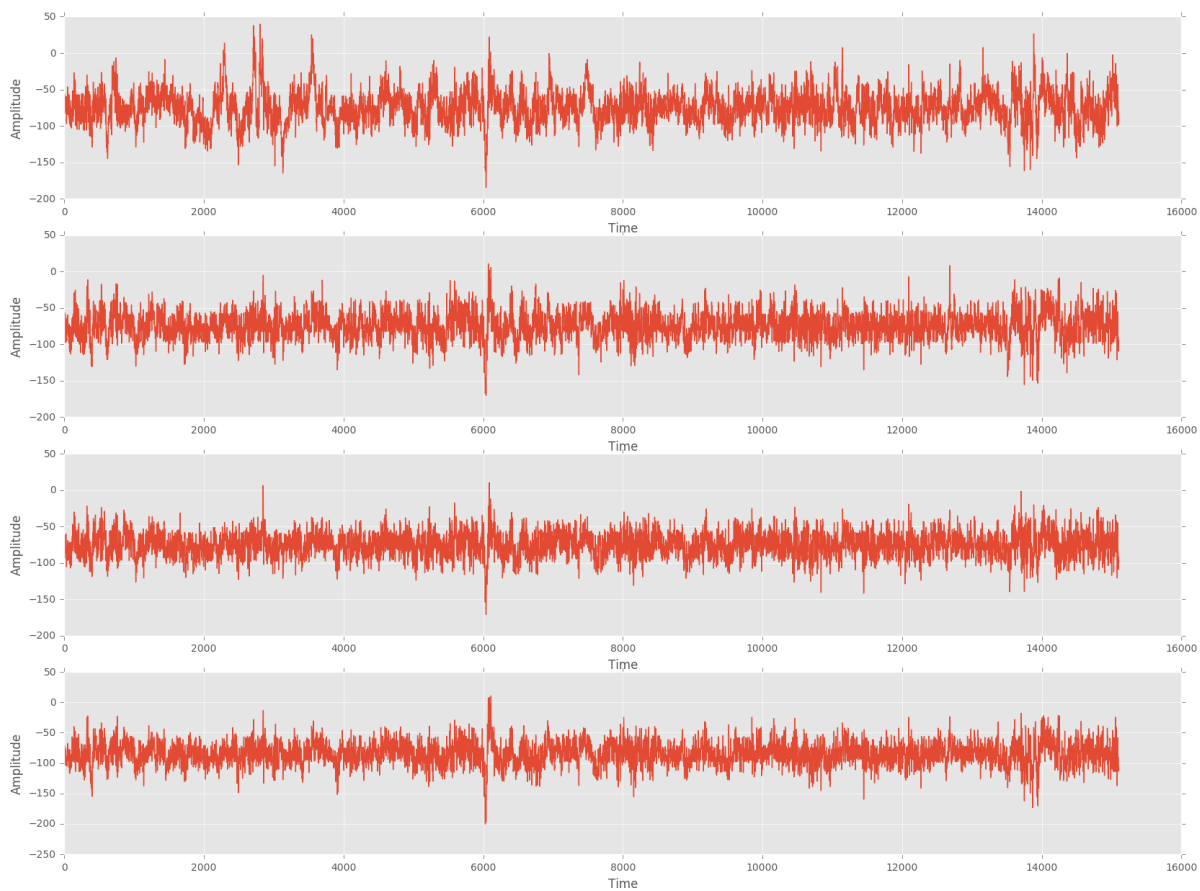
Dataset Generation

The data are in the form of csv files with raw waveform signals from 14 probes places around the scalp. The sampling rate is 128 hz, which allows for frequency analysis up to ~60 hz. Each of 8 subjects participated in two-minute- long sessions. The image below shows the raw waveform data from four of the 4 channels during a typical session. EMG signals (such as those caused by swallowing

or 1 yawning) were manually removed which is shown in figure show below.



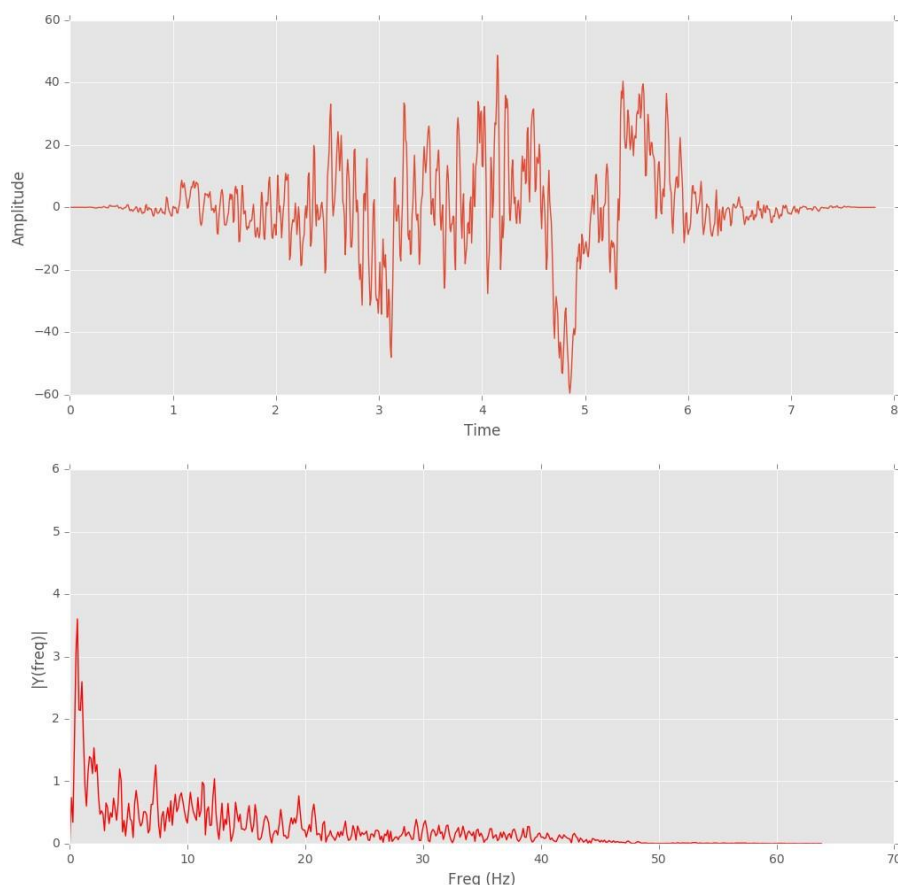
Proposed Block Diagram



Raw waveform data from four of the 14 EEG probes

Hanning Window:

The Hanning window, after its inventor whose name was Von Hann, has the shape of one cycle of a cosine wave with 1 added to it so it is always positive. The sampled signal values are multiplied by the Hanning function, and the result is shown in the figure. Note that the ends of the time record are forced to zero regardless of what the input signal is doing. While the Hanning window does a good job of forcing the ends to zero, it also adds distortion to the wave form being analyzed in the form of amplitude modulation; i.e., the variation in amplitude of the signal over the time record. Amplitude Modulation in a wave form result in sidebands in its spectrum, and in the case of the Hanning window, these sidebands, or side lobes as they are called, effectively reduce the frequency resolution of the analyser by 50%. It is as if the analyser frequency "lines" are made wider. In the illustration here, the curve is the actual filter shape that the FFT analyser with Hanning weighting produces. Each line of the FFT analyser has the shape of this curve -- only one is shown in the figure shown below.



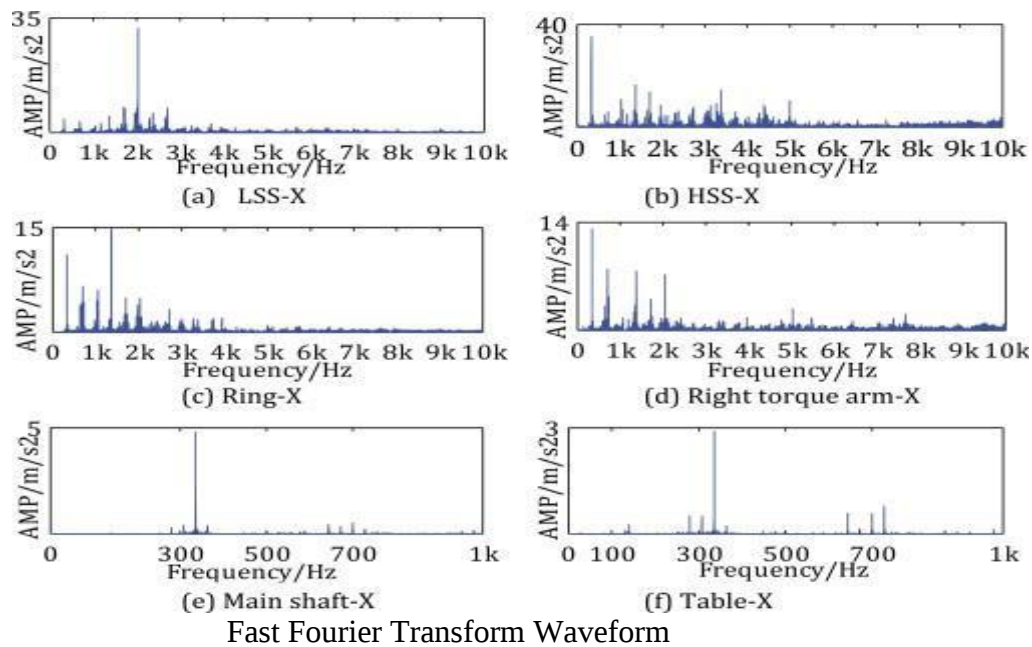
FFT Analyser for Hanning window

If a signal component is at the exact frequency of an FFT line, it will be read at its correct amplitude, but if it is at a frequency that is one half of ΔF (One half the distance between

lines), it will be read at an amplitude that is too low by 1.4 dB. The Hanning window is one out of several attempts to design a window that has favorable properties in the Fourier domain.

Fast Fourier Transform (FFT):

Fast fourier transform (FFT) is one of the most useful tools and is widely used in the signal processing. FFT results of each frame data are listed in figure 6. From figure 6, it can be seen that the vibration frequency are abundant and most of them are less than 5 kHz. Also, the HSS-X point has greater values of amplitude than other points which corresponds with the information provided by the time-domain curve as shown in Figure 4.4.



The fast Fourier transform (FFT), as the name implies, is a fast version of the DFT. The FFT exploits the fact that the straightforward approach to computing the Fourier transform performs the many of the exact same multiplications repeatedly. The FFT algorithm organizes these redundant computations in a very efficient manner by taking advantage of the algebraic properties in the Fourier matrix. Specifically, the FFT makes use of periodicities in the signs that are multiplied to perform the calculation. The result shown in equation 4.2

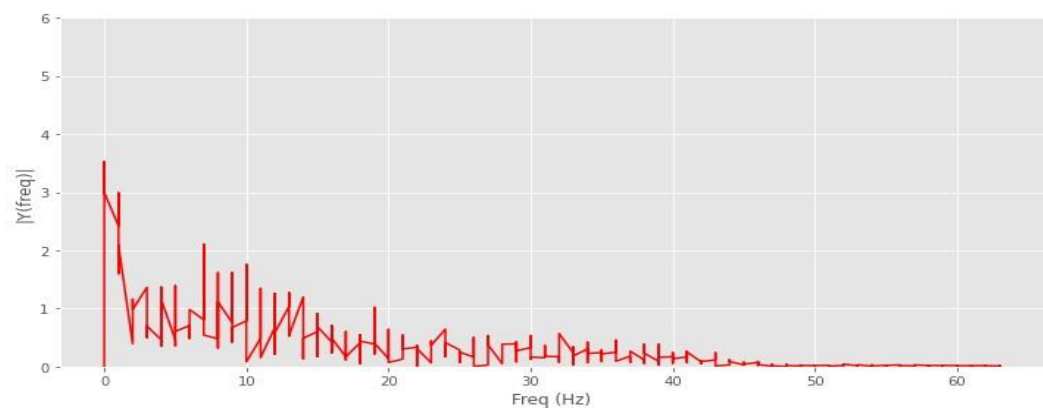
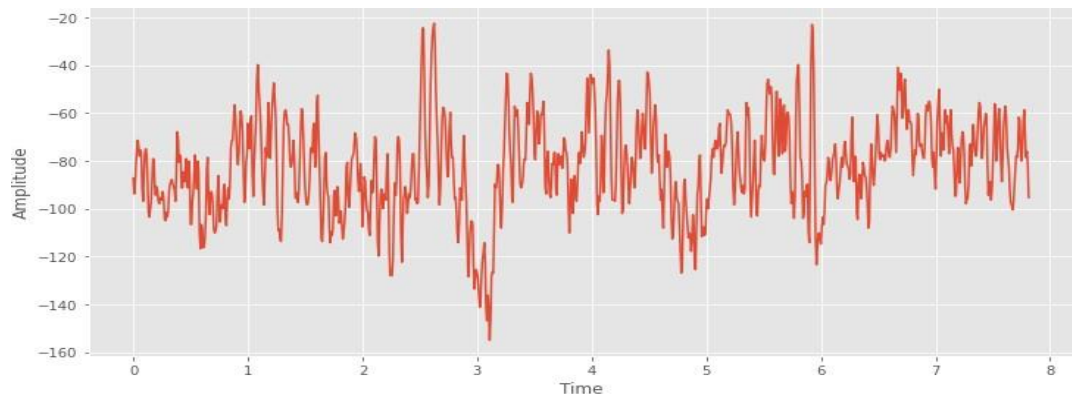
$$X[h] = \sum_{i=0}^{N-1} (x[i] W^{ih}) e^{-j2\pi \frac{h}{N} i}, \quad \text{for } h = 0 \quad (4.2)$$

The frequency spectrum vector is divided into different ranges of frequencies to automate the selection process of the sensitive frequencies to the fault under investigation. The average of each range is then taken as a sensory feature for the system.

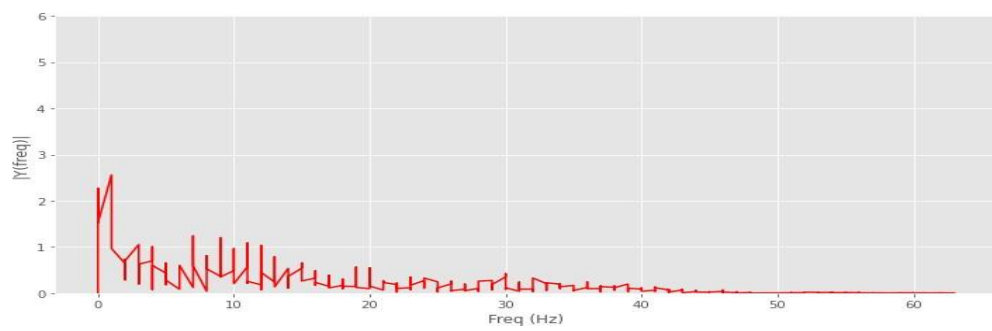
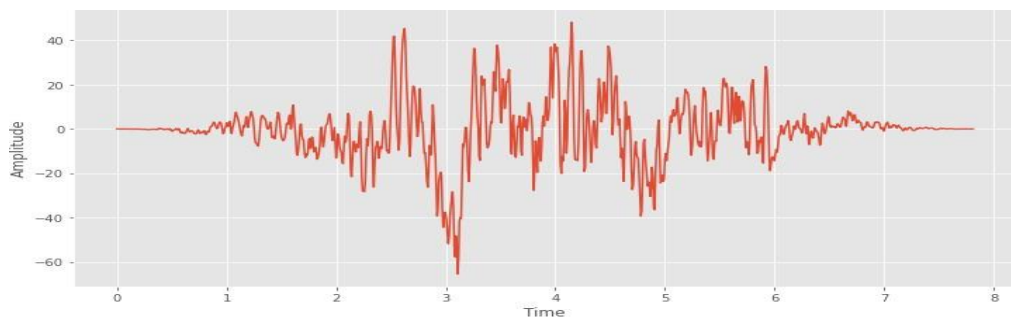
FFT is used in the wind turbine gearbox data analysis, for example for fault detection and evaluation of bearings, shafts, and gears. If damage initiates in a bearing, the shape of the vibration distribution deviates to a gauss-shaped curve. Frequencies are conditioned by the rotational shaft velocity, as well as by shape and size of faults in bearings. Originally, its purpose was served by band selective filters. The advantage of frequency domain analysis over time-domain analysis is its ability to identify and isolate certain frequency component of interest.

Frequency Binning:

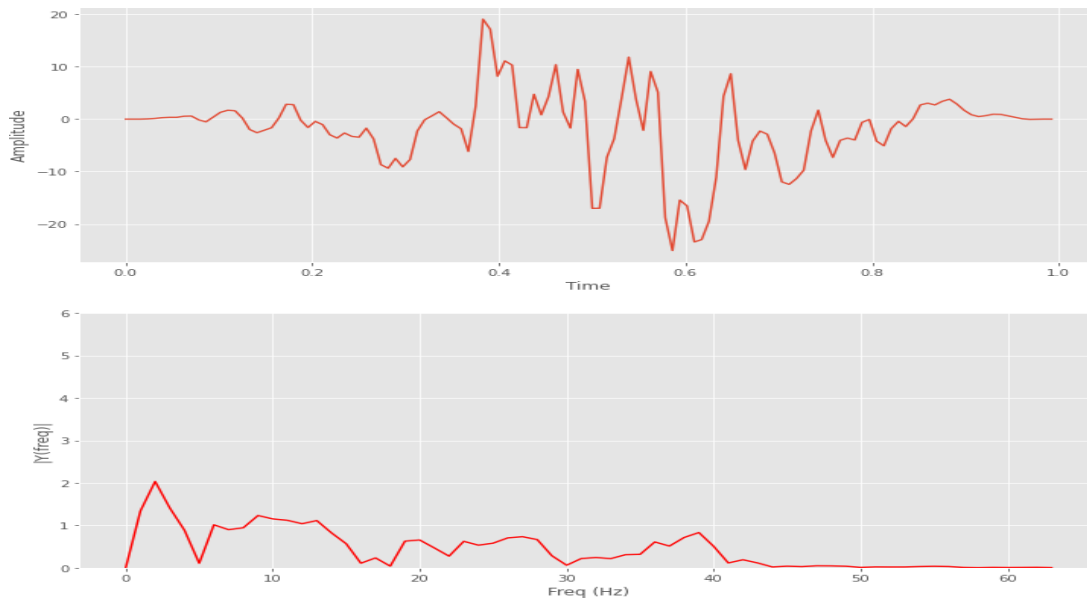
FFT amplitudes were grouped into theta(4-8Hz), alpha(8-12Hz), and beta(12-40Hz) ranges, giving 3 scalar values for each probe per frame. Frequency-dependent Convergence Coefficient apart from the computational advantages, it may also be possible to use the frequency domain implementation to improve the convergence properties of the LMS algorithm. This was originally suggested by Ferrara, who argued that in the frequency domain the error signal in a given frequency bin, $E(k)$, is only a function of the filter coefficient in the same bin, $W(k)$, and so each of the frequency domain filter coefficients converge independently. If this were the case the convergence coefficient could be selected independently for each bin, so that the adaptation algorithm becomes Frequency binning is simple choosing you bin boundaries in a way that the bin content size is the same. For the frequency approach it looks like the order the elements by size and calculate the bin edges in the middle between the highest element of bin A and the lowest of bin B. result as shown in Figures 4.5, 4.6 and 4.7.



Frequency Binning



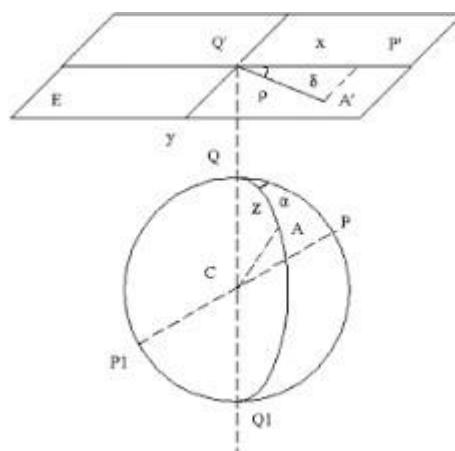
Frequency Binning



Frequency Binning

2D Azimuthal Projection:

A map projection in which a globe, as of the Earth, is assumed to rest on a flat surface onto which its features are projected. An azimuthal projection produces a circular map with a chosen point—the point on the globe that is tangent to the flat surface—at its center. When the central point is either of Earth's poles, parallels appear as concentric circles on the map and meridians as straight lines radiating from the center. Directions from the central point to any other point on the map are accurate, although distances and shapes in some azimuthal projections are distorted away from the center. Compare conic projection cylindrical projection as shown in Figure 4.8.

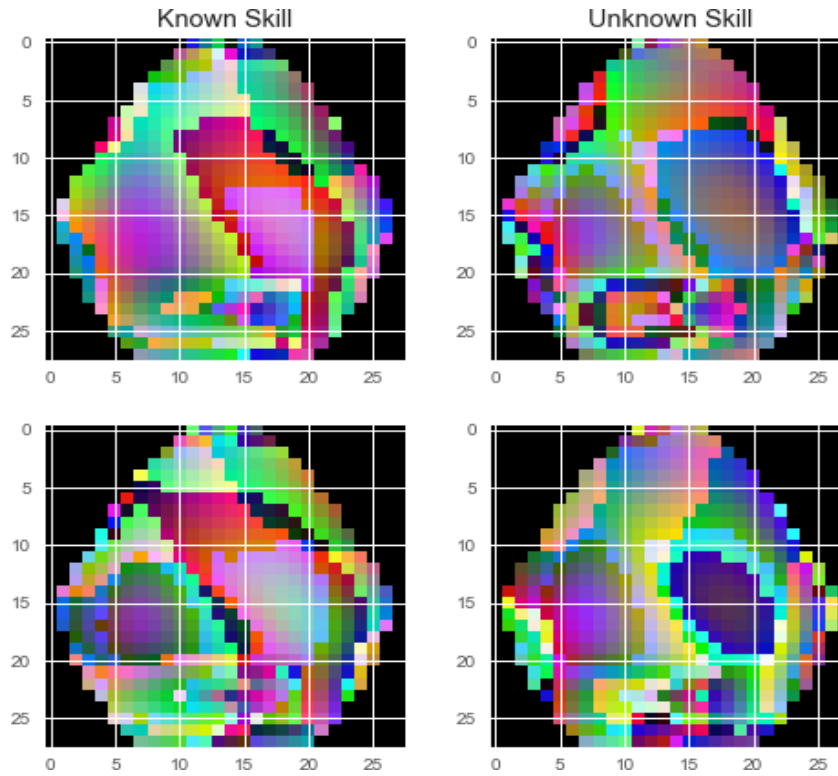


Map projection

Result:

The results obtained are encouraging. Without even using a recurrent neural network the CNN can correctly classify the test subject's brain-state about 8.5 times out of 10. This is likely high enough to enable a new level of performance with brain-computer interface (BCI) technologies. However, the best results were obtained when the network was trained on samples from the same recording session. While this may be practical for basic brain research, it would be less practical for use in BCI technology. The results obtained suggest that while EEG signals do indeed generalize between individuals, there are still significant variations between individuals, which is an unsurprising finding. This further suggests that using EEG for BCI will likely require an iterative approach of training on a large population and then fine tuning on a specific individual. It is therefore recommended that future research be done on the possible application of Transfer Learning techniques to the classification of EEG signals. Training of the dataset is crucial.

We have used the 80% of the dataset for training of the model and 20% of the datasets for testing of the model. We have used only with a smaller number of class samples of persons. We tried to optimize our training process to save computation time and maintain a good overall performance at the same time. We evaluated different kinds of parameter settings and found the following to be very effective. 85% validation accuracy when the CNN had been trained on data from the same EEG session. 81% validation accuracy when then CNN was trained on all individuals but had never seen the test session. 71% validation accuracy when the CNN had been trained on all data from other individuals but had never seen the test individual as shown in Figure 5.1.



Feature projection for known skill and unknown skill

Performance metrics

Performance measure	CNN	Proposed Network
Accuracy	97.57%	92.62%
F-score	94.06%	89.78%
Recall	93.87%	93.26%
Sensitivity	86.79%	82.61%
Specificity	98.52%	89.78%
Precision	94.27%	87.87%

Confusion Matrix:

In this work, the proposed deep learning network was trained on the dataset. Afterwards, the model was evaluated on the test set, which showed good performance. Table 5.1 the performance matrix of the classification results, where each row represents the actual category, while each column stands for the predicted result which is shown in Table below.

Predictions	Classes	
	A	B
Known Skill	230	8
Unknown Skill	4	269

Confusion Matrix

CONCLUSION & FUTURE SCOPE

This work is motivated by the high-level goal of finding robust representations from EEG data, that would be invariant to inter- and intra-subject differences and to inherent noise associated with EEG data collection. We propose a novel methodology for learning representations from multi-channel EEG time-series and demonstrate its advantages in the context of mental load classification task. Our approach is fundamentally different from the previous attempts to learn high- level representations from EEG using deep neural networks. Specifically, rather than representing low-level EEG features as a vector, we transform the data into a sequence of topology-preserving multi-spectral images (EEG” movie”), as opposed to standard EEG analysis techniques that ignore such spatial information. We then train deep recurrent-convolutional networks inspired by state-of-the-art video classification to learn robust representations from the sequence of images. The proposed approach demonstrates significant improvements in classification accuracy over the state-of-the-art results. Since our approach transforms the EEG data into sequence of EEG images, it can be applied on EEG data acquired with different hardware (e.g., with different number of electrodes). The preprocessing step used in our approach transforms the EEG time-series acquired from various sources. Visualization of feature maps and their input activation patterns at various depth levels of convolutional network. The left column (Input EEG images) shows the top 9 images with highest feature activations across the training set. The middle column (Feature Maps) shows the feature map derived in the output of the kernel. Right column (Back Projections) shows the back projected maps derived by applying deconvnet on the feature map displaying structures in the input image that excite that feature map. into comparable EEG frames. In this way, various EEG datasets could be merged. The only information needed to complete this transform would be the spatial coordinates of electrodes for each setup. As a future direction, it would be possible to use unsupervised pretraining methods with larger (or merged) unlabeled EEG datasets prior to training the network with task-specific data.

FACE MASK DETECTION AND ALERTING USING VOICE BY MATLAB

KARTHIKEYAN M, NEWTON PRABU P

ABSTRACT

The main objective of this project is to develop an embedded system, which is used for security applications in face mask . Biometrics technology is rapidly progressing and offers attractive opportunities. In recent years, biometric authentication has grown in popularity as a means of personal identification in ATM authentication systems. The prominent biometric methods that may be used for authentication include fingerprint, palmprint, handprint, face recognition, speech recognition, dental and eye biometrics. Here the project is about the face mask detection and alert using words moreover in lot of public places the people are requested to use mask in the order prevent from the airborne diseases like covid19 etc . Our project identify the people who didn't wear the mask and alert to the authority to take action . Our project can install in many places like schools, colleges, movie theaters etc. We hope by our project we will give awareness to the people to get more caution about this which situation become worse.

1.INTRODUCTION

Biometrics is the science and technology of measuring and analyzing biological data. In information technology, biometrics refers to technologies that measure and analyze human body characteristics, such as DNA, fingerprints, eye retinas and irises, voice patterns, facial patterns and hand measurements, for authentication purposes. In this paper we have used thumb impression for the purpose of face mask identification or authentication. As the thumb impression of every individual is unique, it helps in maximizing the accuracy. A database is created containing the face images of all the face mask in the constituency. Illegal votes and repetition of votes is checked for in this system. Hence if this system is employed the elections would be fair and free from rigging. Thanks to this system that conducting elections would no longer be a tedious and expensive job.

2.WORKING

2.1.FACE MASK RECOGNITION

The goal of this paper is to deal with one class of face recognition problem where some of facial appearances in a given face image are badly deformed by such variations as large expression changes or partial occlusions (or disguise) due to sunglasses, scarves, mustaches and so on. Such variations in facial appearance are commonly encountered in uncontrolled Situations and may cause big trouble to the face-recognition- based security system but are less studied in literatures [5]. Notice that in this paper, we don't intend to deal with other commonly encountered variations in uncontrolled conditions like lighting changes and ageing effect, which are of interest but usually

change people's facial appearance in a more holistic way. By contrast, the facial appearance changes caused by variant expressions and partial occlusions are mostly local in nature, i.e., only parts of facial appearance change largely while others are less affected. The challenge lies in that such local deformations or occlusions in facial appearance can be anywhere and in any size or shape in a given face image and we don't have any prior knowledge about it.

This paper proposes to address this from the aspect of partial similarity matching by exploiting the spatial contiguity nature of occlusions and other local deformations. This idea where many observers will feel that both the person and the horse are similar to the centaur, but the person and the horse are not similar to each other at all. Why? One possible reason is that when comparing two images, human beings tend to "focus on the portions that are very similar and are willing to pay less attention to regions of great dissimilarity" [9]. Inspired by this, one goal of this paper is to design an effective mechanism to support such a robust perception of similarity by humans in face recognition systems, i.e., automatically detecting and capturing the significant partial similarities between two face images while ignoring the unreliable and unimportant features due to expression changes, occlusions or disguises.

The main goal of the pre-processing is to improve the image quality to make it ready to further processing by removing or reducing the unrelated and surplus parts in the background of the mammogram images. Mammograms are medical images that are complicated to interpret. Hence pre-processing is essential to improve the quality. It will prepare the mammogram for the next two-process segmentation and feature extraction. The noise and high frequency components are removed by filters.

Image pre-processing is the name for operations on images at the lowest level of abstraction whose aim is an improvement of the image data that suppresses undesired distortions or enhances some image features important for further processing. It does not increase image information content. Its methods use the considerable redundancy in images. Neighboring pixels corresponding to one object in real images have the same or similar brightness value and if a distorted pixel can be picked out from the image, it can be restored as an average value of neighboring pixels. Image pre-processing tool, created in MatLab, realizes many brightness transformations and local pre-processing methods.

OPERATION

Image pre-processing is the term for operations on images at the lowest level of abstraction. These operations do not increase image information content but they decrease it if entropy is an information measure. The aim of pre-processing is an improvement of the image data that suppresses undesired distortions or enhances some image features relevant for further processing and analysis.

task. Image preprocessing use the redundancy in images. Neighboring pixels corresponding to one real object have the same or similar brightness value. If a distorted pixel can be picked out from the image, it can be restored as an average value of neighboring pixels. Image pre-processing methods can be classified into categories according to the size of the pixel neighborhood that is used for the calculation of a new pixel brightness. In this

IMAGE CROPPING AND FILTERING

The first step in image pre-processing is image cropping. Some irrelevant parts of the image can be removed and the image region of interest is focused. This tool provides a user with the size information of the cropped image. MatLab function for image cropping realizes this operation interactively waiting for an user to specify the crop rectangle with the mouse and operates on the current axes. The output image is of the same class as the input image. The two-dimensional convolution operation is fundamental to the analysis of images. A new value is ascribed to a given pixel based on the evaluation of a weighted average of pixel values in a $k \times k$ neighborhood of the central pixel. Convolution kernel or the filter mask is represented with weights supplied in a square matrix. It is applied to each pixel in an image. Discrete form of the 2D convolution operator is defined by the following relationship between the elements (x, y) of the input image, the elements $h(\alpha, \beta)$ of the convolution kernel, and the elements $g(x, y)$ of the output image by the following master formula

$$g(x, y) = \sum_{\alpha=-(k-1)/2}^{(k-1)/2} \sum_{\beta=-(k-1)/2}^{(k-1)/2} f_i(\alpha, \beta) h(x - \alpha, y - \beta),$$

x, y, α and β are integers. Coefficients of the kernel H represent a discrete approximation of the analytical form of the response function characterizing the desired filter. In practical cases, the kernel is a square array and $k_x = k_y = k$, where k is odd and much smaller than the linear image dimension. There is the following steps, realized for each pixel P represented by (x, y) :

- placement of H on P
- multiplication of each pixel in the $k \times k$ neighborhood by the appropriate filter mask
- summation of all products
- placement of the normalized sum into position P of the output image.

This tool for pre-processing lets an user explores 2-D Finite Impulse Response filters. By changing the cut-off frequency and filter order, the user can design filter and can see the designed filter's coefficients and frequency response. Median filtering is a non-linear smoothing method that reduces the blurring of edges and significantly eliminates impulse noise. It suppresses image noise without reducing the image sharpness and can be applied iteratively. The brightness value of the

current pixel in the image is replaced by the median brightness of either 3-by-3 or 4-by-4 neighborhood.

INTENSITY ADJUSTEMENT AND HISTOGRAM EQUALIZATION

A gray-scale transformation T of the original brightness p from scale $[p_0, p_k]$ into brightness q from a new scale $[q_0, q_k]$ is given by $q = T(p)$. It does not depend on the position of the pixel in the image. Values below p_0 and above p_k are clipped. Values below p_0 map to q_0 , and those above p_k map to q_k . Alpha argument specifies the shape of the curve describing the relationship between the values in the input image and output image. If alpha is less than 1, the mapping is weighted toward brighter output values. If alpha is greater than 1, the mapping is weighted toward lower darker output values. If the argument is omitted its default value is 1. Graphical controls enable an user to increase and decrease the brightness, contrast and alpha correction.

Another offered possibility to enhance the contrast of image, by transforming the values in an intensity image so that the histogram of the output image matches a specified histogram, is histogram equalization technique. Region description is based on its statistical gray-level properties. Histogram provides the frequency of the brightness value in the image. An image with n gray levels is represented with one-dimensional array with n elements. The n -th element of array contains the number of pixels whose gray level is n .

Assume that the pixel values are normalized and lie in the range $[0, 1]$. Let $s = T(r)$, for any $r \in [0, 1]$, is transformation function which satisfies the following conditions:

- $T(r)$ is single valued and monotonically increasing in the interval $[0, 1]$;
- $0 \leq T(r) \leq 1$ for any $r \in [0, 1]$.

The original and transformed gray levels can be characterized by their probability density functions. Contrast is the local change in brightness and is defined as the ratio between average brightness of an object and the background brightness. Histogram equalization technique is based on modifying the appearance of an image by controlling the probability density function of its gray levels by the transformation function $T(r)$. This technique enhances the contrast in the image.

BRIGHTNESS THRESHOLDING

Brightness thresholding is an indispensable step in extracting pertinent information. A gray-scale image often contains only two level of significant information: the foreground level constituting objects of interest and the background level against which the foreground is discriminated [1]. A complete segmentation of an image R is a finite set of regions $R_1, R_2, \dots, R_m$,

$$R = \bigcup_{i=1}^m R_i, \quad R_i \cap R_j = \emptyset \quad i \neq j.$$

If R_b is a background in the image, then $\bigcup_{i=1}^m R_i$ is considered the object and $R_c = \bigcup_{i=1}^m R_i$, where R_c is the set complement. While there are two principal peaks of the foreground and the background intensities in the image histogram, there are many other gray intensities present. Binarization can be accomplished by the choice of an intensity, between the two histogram peaks, that is the threshold between all background intensities below and all foreground intensities above. The input image I_1 is being transformed to an output binary segmented image I_2 , in the following way

$$I_2(i, j) = \begin{cases} 1; & I_1(i, j) \geq T \\ 0; & I_1(i, j) < T \end{cases}$$

where T is the threshold. $I_2(i, j) = 1$ for the object elements and $I_2(i, j) = 0$ for the background elements. There are different approaches for image binarization depending on the type of image. Successful threshold segmentation depends on the threshold selection.

A number of conditions like poor image contrast or spatial non uniformities in background intensity can make difficult to resolve foreground from background. These cases require user interaction for specifying the desired object and its distinguishing intensity features.

CLEARRING AREAS OF A BINARY IMAGE

If there is a deformation of the expected shape and size of the border and the whole region during the separation of image object from its background, it can be partially overcome. Usually small polygon mask, located next to the region and out of it, is added to clear image area with similar brightness of the region. This mask can reshape image objects and provides a separation image objects from each other and from their image background. This operation is realized interactively, adding vertices to the polygon. Selecting a final vertex of the polygon over a white colored image region, the fill is started and recolor to black. Created fill is a logical mask and the input image is logical matrix. Using logical operator AND under logical arguments, the output image is also obtained as a logical array. If the white regions represents image foreground on the black background, whole objects or their parts can be deleted. Special user requirements about the size and shape of the observed object, can be realized by the same way.

DETECTING EDGES

Edges are pixels where the intensity image function changes abruptly. Edge detectors are collection of local image pre-processing methods used to locate changes in the brightness function. An image function depends on two variables, co-ordinates in the image plane. Operators describing edges are expressed by partial derivatives. A change of the image function can be described by a gradient that points in the direction of the largest growth of the image function.

An edge is a vector variable with two components, magnitude and direction. The edge magnitude is the magnitude of the gradient. The edge direction is rotated with respect to the gradient direction by $-\pi/2$. The gradient direction gives the direction of maximum growth of the function, e.g., from black to white. The boundary and its parts are perpendicular to the direction of the gradient. The gradient magnitude and gradient direction

$$|grad\ g(x, y)| = \sqrt{\left(\frac{\partial g}{\partial x}\right)^2 + \left(\frac{\partial g}{\partial y}\right)^2}$$

$$\varphi = arg\left(\frac{\partial g}{\partial x}, \frac{\partial g}{\partial y}\right)$$

are continuous image functions where $arg(x, y)$ is the angle from x axis to the point (x, y).

A digital image is discrete in nature and these equations must be approximated by differences of the image g, in the vertical direction for fixed i and in the horizontal direction for fixed j, by following equations

$$\begin{aligned}\Delta_i g(i, j) &= g(i, j) - g(i - n, j) \\ \Delta_j g(i, j) &= g(i, j) - g(i, j - n),\end{aligned}$$

where n is a small integer chosen so to provide a good approximation to the derivative and to neglect unimportant changes in the image function. Gradient operators approximating derivatives of the image function using differences use one or several convolution masks. Beside them, this tool uses operators based on the zero-crossings of the image function second derivative. Sobel, Prewitt, and Roberts methods find edges by thresholding the gradient and by them horizontal edges, vertical edges or both can be detected.

The Laplacian of Gaussian method thresholds the slope of the zero crossings after filtering the image with a LoG filter [3]. Canny method thresholds the gradient using the derivative of a Gaussian filter. One option in the main menu provides processing of the image region of interest or the whole image. It has to move the pointer over the image on the left and when the cursor changes to a crosshair, it has to click on points in the image to select vertices of the region of interest. Offered operations to

In the present days it is being used for computer network access and entry devices for building door locks. Face Mask recognition are being used by banks for authorization and are becoming more common at grocery stores where they are utilized to automatically recognize a registered customer and bill their credit card or debit account. Finger-scanning technology is being used in a novel way at some places where cafeteria purchases are supported by a federal subsidized meal program. ted.

HANDWRITTEN DIGIT RECOGNITION USING MACHINE LEARNING

E.SALMANKHAN , V.DHARANIDHARAN

ABSTRACT

Digit Recognition is a noteworthy and important issue. As the manually written digits are not of a similar size, thickness, position and direction, in this manner, various difficulties must be considered to determine the issue of handwritten digit recognition. The uniqueness and assortment in the composition styles of various individuals additionally influence the example and presence of the digits. It is the strategy for perceiving and arranging transcribed digits. It has a wide range of applications, for example, programmed bank checks, postal locations and tax documents and so on. In this project we have to create an GUI by using the python programming language. We have to implement the Tkinter python library to create an GUI interface for predicting the digit and the Tkinter is denoted as tk. On that app contains two buttons one is clear and another one is recognize. The recognize button is used to recognize the digit what we are drawing on the screen and clear button is used to clear the digit what we draw in the last and again we have to draw a new digit to recognize continuously. Right side of the app we have to see the number what we draw and it shows the accuracy of the number also. We have to build and train the convolutional neural network which is very effective for image classification purposes. Later on, we build an GUI where we are draw the digit on the canvas then we classify the digit and show the results.

1.INTRODUCTION

The interesting Python project requires you to have basic knowledge of Python programming, deep learning with Keras library and the Tkinter library for building GUI. Install the necessary libraries for this project using this command: `pip install numpy`

This is probably one of the most popular datasets among machine learning and deep learning enthusiasts. The [MNIST dataset](#) contains 60,000 training images of handwritten digits from zero to nine and 10,000 images for testing. So, the MNIST dataset has 10 different classes. The handwritten digits images are represented as a 28×28 matrix where each cell contains grayscale pixel value.

Data selection is defined as the process of determining the appropriate data type and source, as well as suitable instruments to collect data. Data selection precedes the actual practice of data collection. This definition distinguishes data selection from selective data reporting (selectively excluding data that is not supportive of a research hypothesis) and interactive/active data selection (using collected data for monitoring activities/events, or

conducting secondary data analysis). The process of selecting suitable data for a research project can impact data integrity.

The primary objective of data selection is the determination of appropriate data type, source, and instrument(s) that allow investigators to adequately answer research questions. This determination is often discipline-specific and is primarily driven by the nature of the investigation, existing literature, and accessibility to necessary data sources.

Integrity issues can arise when the decisions to select 'appropriate' data to collect are based primarily on cost and convenience considerations rather than the ability of data to adequately answer research questions. Certainly, cost and convenience are valid factors in the decision-making process. However, researchers should assess to what degree these factors might compromise the integrity of the research endeavor.

There are a number of issues that researchers should be aware of when selecting data. These include determining:

- the appropriate type and sources of data which permit investigators to adequately answer the stated research questions,
- suitable procedures in order to obtain a **representative sample**
- the proper instruments to collect data. There should be compatibility between the type/**source** of data and the mechanisms to collect it. It is difficult to extricate the selection of the type/source of data from instruments used to collect the data.

First, we are going to import all the modules that we are going to need for training our model. The Keras library already contains some datasets and MNIST is one of them. So we can easily import the dataset and start working with it. The **mnist.load_data()** method returns us the training data, its labels and also the testing data and its labels.

The image data cannot be fed directly into the model so we need to **perform some operations and process the data** to make it ready for our neural network. The dimension of the training data is (60000,28,28). The CNN model will require one more dimension so we reshape the matrix to shape (60000,28,28,1).

Now we will **create our CNN model** in Python data science project. A CNN model generally consists of convolutional and pooling layers. It works better for data that are represented as grid structures, this is the reason why CNN works well for image classification problems. The dropout layer is used to deactivate some of the neurons and while training, it

reduces overfitting of the model. We will then compile the model with the Adadelta optimizer.

The **model.fit()** function of Keras will start the training of the model. It **takes the training data, validation data, epochs, and batch size**. It takes some time to train the model. After training, we save the weights and model definition in the 'mnist.h5' file.

We have 10,000 images in our dataset which will be used to **evaluate how good our model works**. The testing data was not involved in the training of the data therefore, it is new data for our model. The MNIST dataset is well balanced so we can get around 99% accuracy. Now for the GUI, we have created a new file in which we **build an interactive window to draw digits on canvas** and with a button, we can recognize the digit. The Tkinter library comes in the Python standard library. We have created a function **predict_digit()** that takes the image as input and then uses the trained model to predict the digit.

Then we **create the App class** which is responsible for building the GUI for our app. We create a canvas where we can draw by capturing the mouse event and with a button, we trigger the **predict_digit()** function and display the results.

2.OBJECTIVES AND PROBLEM STATEMENT

The aim of a handwriting digit recognition system is **to convert handwritten digits into machine readable formats**. The main objective of this work is to ensure effective and reliable approaches for recognition of handwritten digits and make banking operations easier and error free.

Recently handwritten digit recognition becomes vital scope and it is appealing many researchers because of its using in variety of **machine learning and computer vision applications**. However, there are deficient works accomplished on Arabic pattern digits because Arabic digits are more challenging than English patterns.

Handwritten digit recognition is the ability of a computer to recognize the human handwritten digits from different sources like images, papers, touch screens, etc, and classify them into 10 predefined classes (0-9). This has been a topic of boundless-research in the field of deep learning.

2.1 Problem Statement

The **handwritten digit recognition** is the capability of computer applications to **recognize** the human **handwritten digits**. It is a hard task for the **machine** because **handwritten digits** are not perfect and can be made with many different shapes and sizes. The **handwritten digit recognition system** is a way to tackle this problem which uses

the image of a **digit** and recognizes the **digit** present in the image. Convolutional **Neural Network** model created using **PyTorch library** over the **MNIST dataset** to **recognize handwritten digits**. To identify the requirement of Convolutional Neural Network for handwritten recognition. To evaluate the accuracy of Convolutional Neural Network model for classify the handwritten digits. Create an app for predict the handwritten digits so that we can recognize the digits easily.

We have to develop the proposed network using training Samples in MNIST database. Achieving the good test result by using this method. Training the dataset is place a major role in this project because sometimes the training data contains some error on that time what we do we have to modify the program based upon our reuirement and again test after we get a good result.

2.2 Problem Definition

The handwritten digit recognition is **the capability of computer applications to recognize the human handwritten digits**. It is a hard task for the machine because handwritten digits are not perfect and can be made with many different shapes and sizes. The handwritten digit recognition system is a way to tackle this. Handwritten digit recognition is the ability of a computer to recognize the human handwritten digits from different sources like images, papers, touch screens, etc, and classify them into 10 predefined classes (0-9). The different stages of Handwritten character recognition system are **Pre-processing, Segmentation, Feature Extraction and Classification**. These are the steps involved in this project and there are other some steps are also available that is importing the necessar libraries, Create a model, Train the model, Evaluate the model, create a GUI to predict the handwritten digits.

A neural network is a model inspired by how the brain works. It consists of multiple layers having many activations, this activation resembles neurons of our brain. A neural network tries to learn a set of parameters in a set of data which could help to recognize the underlying relationships. Neural networks can adapt to changing input; so the network generates the best possible result without needing to redesign the output criteria.

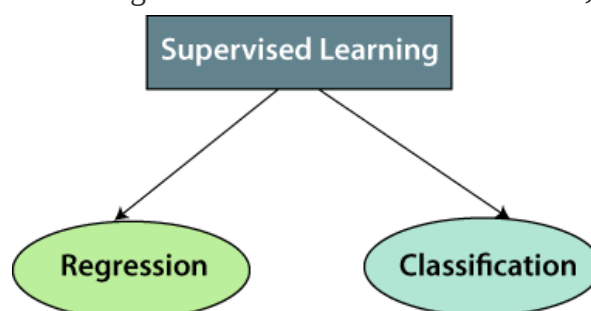
3. METHODOLOGY

Machine learning (ML) is **a type of artificial intelligence (AI) that allows software applications to become more accurate at predicting outcomes without being explicitly programmed to do so**. Machine learning algorithms use historical data as input to predict new output values. **Machine learning (ML)** is a field of inquiry devoted to understanding and building methods that 'learn', that is, methods that leverage data to improve

performance on some set of tasks.[1] It is seen as a part of artificial intelligence. Machine learning algorithms build a model based on sample data, known as training data, in order to make predictions or decisions without being explicitly programmed to do so.[2] Machine learning algorithms are used in a wide variety of applications, such as in medicine, email filtering, speech recognition, and computer vision, where it is difficult or unfeasible to develop conventional algorithms to perform the needed tasks.[3]. A subset of machine learning is closely related to computational statistics, which focuses on making predictions using computers, but not all machine learning is statistical learning. The study of mathematical optimization delivers methods, theory and application domains to the field of machine learning. Data mining is a related field of study, focusing on exploratory data analysis through unsupervised learning.[5][6] Some implementations of machine learning use data and neural networks in a way that mimics the working of a biological brain.[7][8] In its application across business problems, machine learning is also referred to as predictive analytics.

3.1 Supervised Machine Learning

Supervised learning is the types of machine learning in which machines are trained using well "labelled" training data, and on basis of that data, machines predict the output. The labelled data means some input data is already tagged with the correct output. In supervised learning, the training data provided to the machines work as the supervisor that teaches the machines to predict the output correctly. It applies the same concept as a student learns in the supervision of the teacher. Supervised learning is a process of providing input data as well as correct output data to the machine learning model. The aim of a supervised learning algorithm is to **find a mapping function to map the input variable(x) with the output variable(y)**. In the real-world, supervised learning can be used for **Risk Assessment, Image classification, Fraud**



Detection, spam filtering, etc.

3.2 Unsupervised Machine Learning

As the name suggests, unsupervised learning is a machine learning technique in which models are not supervised using training dataset. Instead, models itself find the hidden patterns and insights from the given data. It can be compared to learning which takes place in the human

brain while learning new things. It can be defined as:

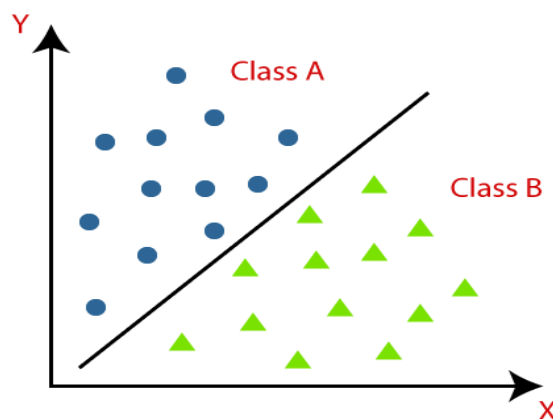
Unsupervised learning is a type of machine learning in which models are trained using unlabeled dataset and are allowed to act on that data without any supervision.

Unsupervised learning cannot be directly applied to a regression or classification problem because unlike supervised learning, we have the input data but no corresponding output data.

The goal of unsupervised learning is to **find the underlying structure of dataset, group that data according to similarities, and represent that dataset in a compressed format**. Suppose the unsupervised learning algorithm is given an input dataset containing images of different types of cats and dogs. The algorithm is never trained upon the given dataset, which means it does not have any idea about the features of the dataset. The task of the unsupervised learning algorithm is to identify the image features on their own. Unsupervised learning algorithm will perform this task by clustering the image dataset into the groups according to similarities between images.

4.1 Implementation of Classification Algorithm

The Classification algorithm is a Supervised Learning technique that is used to identify the category of new observations on the basis of training data. In Classification, a program learns from the given dataset or observations and then classifies new observation into a number of classes or groups. Such as, **Yes or No, 0 or 1, Spam or Not Spam, cat or dog**, etc. Classes can be called as targets/labels or categories. Unlike regression, the output variable of Classification is a



category, not a value, such as "Green or Blue", "fruit or animal", etc. Since the Classification algorithm is a Supervised learning technique, hence it takes labeled input data, which means it contains input with the corresponding output. In classification algorithm, a discrete output function(y) is mapped to input variable(x). $y=f(x)$, where y = categorical output. The best example of an ML classification algorithm is **Email Spam Detector**. The main goal of the Classification algorithm is to identify the category of a given dataset, and these algorithms are mainly used to predict the output for the categorical data. Classification algorithms can be better understood using the below diagram. In the below diagram, there are two classes, class A and Class B. These classes

have features that are similar to each other and dissimilar to other classes.

Random Forest

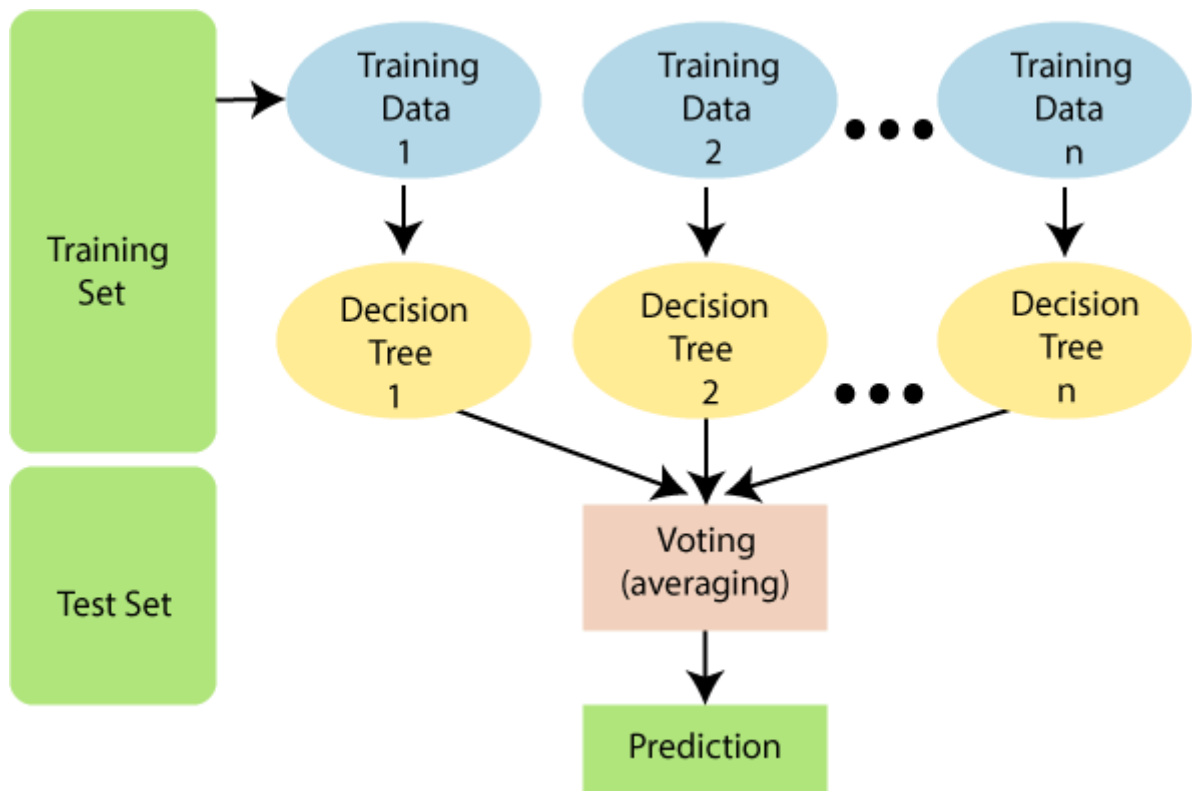
Random Forest is a popular machine learning algorithm that belongs to the supervised learning technique. It can be used for both Classification and Regression problems in ML. It is based on the concept of **ensemble learning**, which is a process of combining multiple classifiers to solve a complex problem and to improve the performance of the model.

As the name suggests, "**Random Forest is a classifier that contains a number of decision trees on various subsets of the given dataset and takes the average to improve the predictive accuracy of that dataset.**" Instead of relying on one decision tree, the random forest takes the prediction from each tree and based on the majority votes of predictions, and it predicts the final output. The greater number of trees in the forest leads to higher accuracy and prevents the problem of over fitting. The below diagram explains the working of the Random Forest algorithm:

Decision Tree Algorithm

- Decision Tree is a **Supervised learning technique** that can be used for both classification and Regression problems, but mostly it is preferred for solving Classification problems. It is a tree- structured classifier, where **internal nodes represent the features of a dataset, branches represent the decision rules** and **each leaf node represents the outcome.**
- In a Decision tree, there are two nodes, which are the **Decision Node** and **Leaf Node**. Decision nodes are used to make any decision and have multiple branches, whereas Leaf nodes are the output of those decisions and do not contain any further branches.
- The decisions or the test are performed on the basis of features of the given dataset.

- It is a graphical representation for getting all the possible solutions to a problem/decision based on given conditions.



- It is called a decision tree because, similar to a tree, it starts with the root node, which expands on further branches and constructs a tree-like structure.
- In order to build a tree, we use the **CART algorithm**, which stands for **Classification and Regression Tree algorithm**.
- A decision tree simply asks a question, and based on the answer (Yes/No), it further split the tree into subtrees.
- Decision Trees usually mimic human thinking ability while making a decision, so it is easy to understand.
- The logic behind the decision tree can be easily understood because it shows a tree-like structure.

In a decision tree, for predicting the class of the given dataset, the algorithm starts from the root node of the tree. This algorithm compares the values of root attribute with the record (real dataset) attribute and, based on the comparison, follows the branch and jumps to the next node.

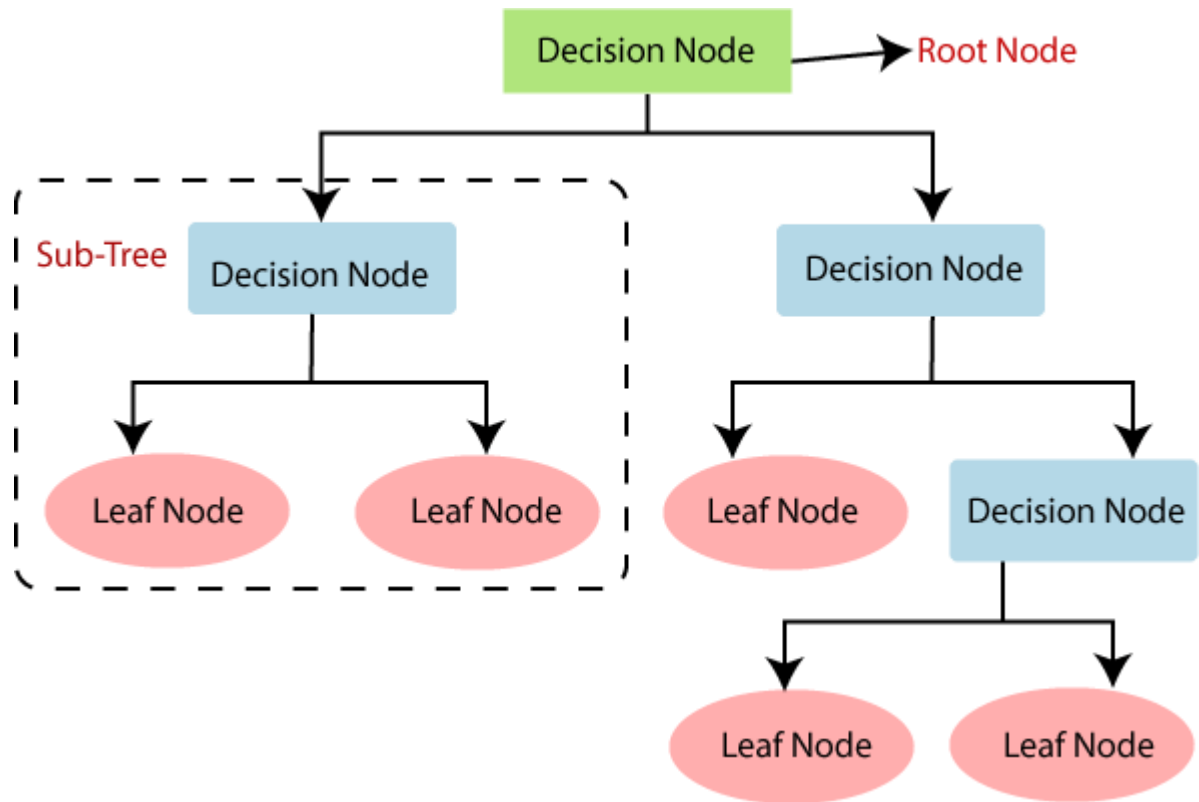
- Below diagram explains the general structure of a decision tree:

Decision Tree Terminologies

- **Root Node:** Root node is from where the decision tree starts. It represents the entire dataset, which further gets divided into two or more homogeneous sets.
- **Leaf Node:** Leaf nodes are the final output node, and the tree cannot be segregated further after getting a leaf node.
- **Splitting:** Splitting is the process of dividing the decision node/root node into sub-nodes according to the given conditions.
- **Branch/Sub Tree:** A tree formed by splitting the tree.
- **Pruning:** Pruning is the process of removing the unwanted branches from the tree.
- **Parent/Child node:** The root node of the tree is called the parent node, and other nodes are called the child nodes.

4. EXISTING SYSTEM AND PROPOSED SYSTEM

The neural networks are widely used for the recognition of characters [1, 2, 3, 4, 5, 6, 7, and 8]. In this work we train and test a neural network classifier using MNIST database, the important step in the recognition is learning, we used descent of the gradient algorithm. In the training process the synaptic weight of the connections between the neurons is modified. The first layers contain attached binary neurons without editable connections. We have proposed the MLP as a classifier used for the recognition of the binary images (black and white pixels). The MLP contains three layers of neurons, the first layer corresponds to the retina. In technical terms it matches the input image. The second layer (hidden layer) corresponds to the extraction of characteristics subsystems. The third layer corresponds to the output system. Each neuron in this layer corresponds to one of the output classes. In the recognition task of handwritten digits, this layer contains 10 neurons corresponding to the digits 0 ... 9 (Figure 1). The original weights of the network at random connections. The weights are changed during the formation of the perceptron. The rule of change of weight corresponds to the learning algorithm. The neural networks are widely used for the recognition of characters [1, 2, 3, 4, 5, 6, 7, and 8]. In this work we train and test a neural network classifier using MNIST database, the important step in the recognition is learning,



We used descent of the gradient algorithm. In the training process the synaptic weight of the connections between the neurons is modified. The first layers contain attached binary neurons without editable connections. We have proposed the MLP as a classifier used for the recognition of the binary images (black and white pixels). The MLP contains three layers of neurons, the first layer corresponds to the retina. In technical terms it matches the input image. The second layer (hidden layer) corresponds to the extraction of characteristics subsystems. The third layer corresponds to the output system. Each neuron in this layer corresponds to one of the output classes. In the recognition task of handwritten digits, this layer contains 10 neurons corresponding to the digits 0 ... 9 (Figure 1). The original weights of the network at random connections. The weights are changed during the formation of the perceptron. The rule of change of weight corresponds to the learning algorithm.

Proposed System

The proposed handwritten digit recognition system follows the standard model of feature based classification systems consisting of the digit image database, an essential feature extraction sub-block and a main classification sub-block. The MNIST Benchmark database of handwritten digits has been considered in this work. The Histogram of Oriented Gradient (HOG) technique has been extended in this work by using a new Multiple-Cell Size (MCS) HOG approach to extract features from the database images and the classification sub-block is based on the Support Vector Machine (SVM) classification methodology. The performance of the classification system depends on a number of factors including the type of feature extraction

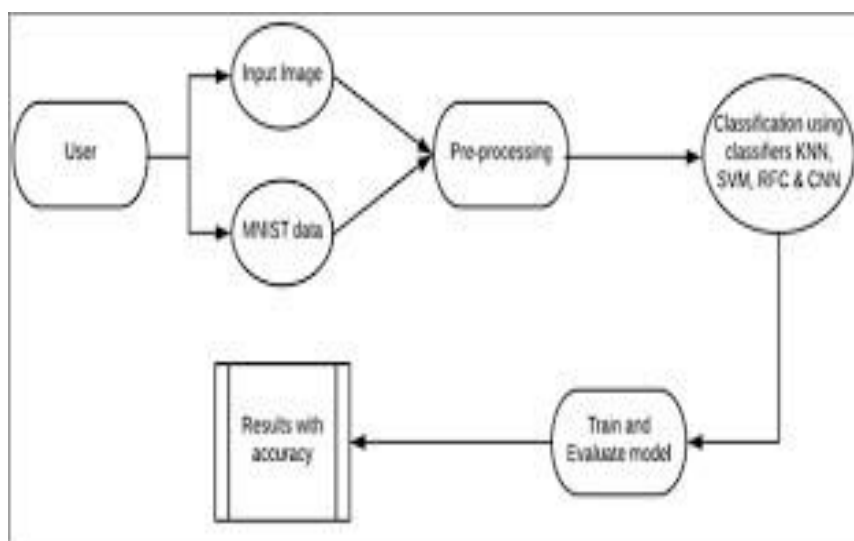
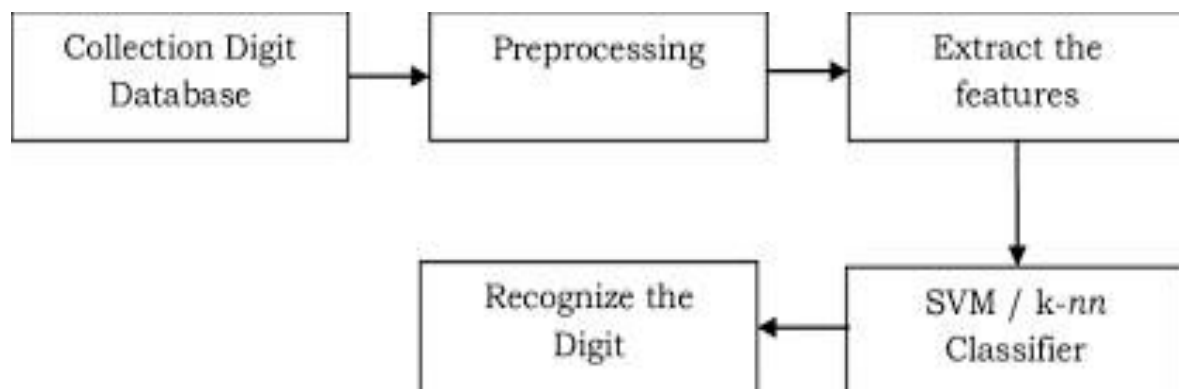
techniques used and the kind of classifiers used for the classification task. Other important factors that affect the classification performance include pre and post processing procedures. In this work the pre or post processing steps have not been employed in the feature extraction process since the HOG descriptors based classification systems are not sensitive to pre-processing operations. The number of images available in the dataset for training, validation and test purposes also plays dominant role in the performance of the classification system. Both the Independent Test Set as well as the 10-Fold Cross-Validation strategies have been used to evaluate the performance of the system. The Block Diagram of the proposed Handwritten Digit Recognition system is shown in the detailed description of the various sub-blocks is given below.

Activity Diagram

An activity diagram visually presents **a series of actions or flow of control in a system** similar to a flowchart or a data flow diagram. Activity diagrams are often used in business process modeling. They can also describe the steps in a use case diagram. Activities modeled can be sequential and concurrent. The basic purposes of activity diagrams is similar to other four diagrams. It captures the dynamic behavior of the system. Other four diagrams are used to show the message flow from one object to another but activity diagram is used to show message flow from one activity to another. Activity is a particular operation of the system. Activity diagrams are not only used for visualizing the dynamic nature of a system, but they are also used to construct the executable system by using forward and reverse engineering techniques. The only missing thing in the activity diagram is the message part.

It does not show any message flow from one activity to another. Activity diagram is sometimes considered as the flowchart. Although the diagrams look like a flowchart, they are not. It shows different flows such as parallel, branched, concurrent, and single. The purpose of an activity diagram can be described as Draw the activity flow of a system. Describe the sequence from one activity to another. Describe the parallel, branched and concurrent flow of the system.

Before drawing an activity diagram, we should identify the following elements – Activities
, Association Conditions Constraints



Module Description

This course provides an easy-to-understand overview of machine learning for anyone interested in how it works, what it can and cannot do and how it is commonly utilized in support of business goals. The course covers common algorithm types and further explains how machine learning systems work behind the scenes. This module provides the basis for the rest of the course by introducing the basic concepts behind machine learning, and, specifically, how to perform machine learning by using Python and the scikit-learn machine learning module. First, you will learn about the basic types of machine learning. Next, you will learn an important step before applying machine learning algorithms, data pre-processing. Finally, you will learn how to leverage different types of machine learning algorithms in a Python script.

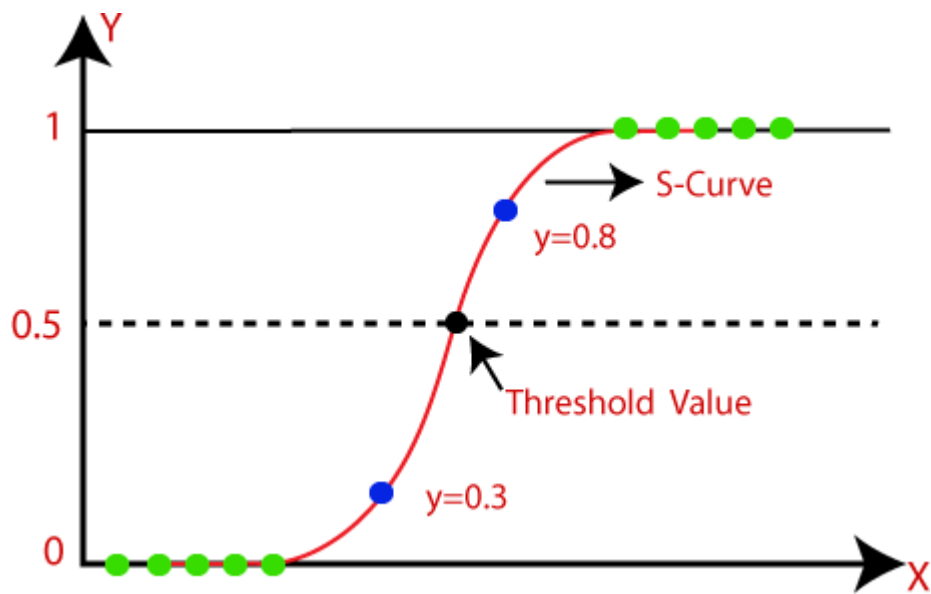
This course, Machine Learning for Accounting with Python, introduces machine learning algorithms (models) and their applications in accounting problems. It covers

classification, regression, clustering, text analysis, time series analysis. It also discusses model evaluation and model optimization. This course provides an entry point for students to be able to apply proper machine learning models on business related datasets with Python to solve various problems. Accounting Data Analytics with Python is a prerequisite for this course. This course is running on the same platform (Jupyter Notebook) as that of the prerequisite course. While Accounting Data Analytics with Python covers data understanding and data preparation in the data analytics process, this course covers the next two steps in the process, modeling and model evaluation. Upon completion of the two courses, students should be able to complete an entire data analytics process with Python.

Logistic Regression

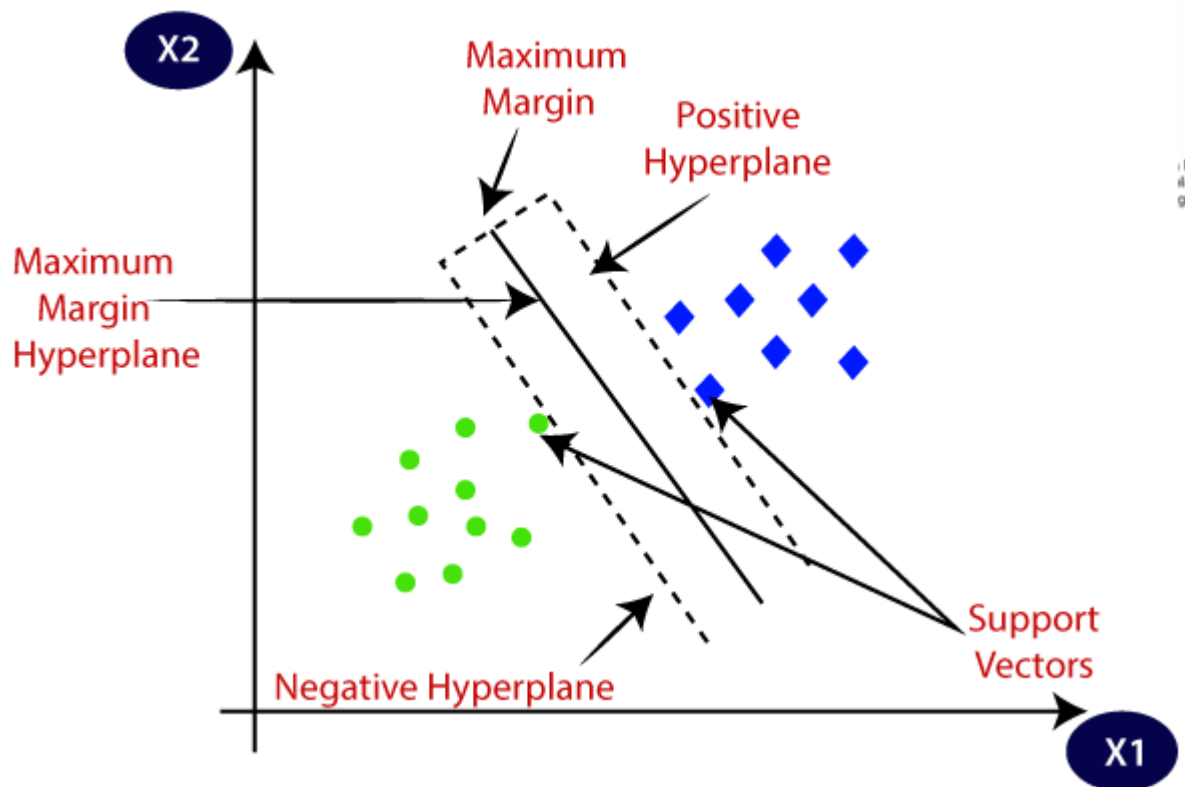
- Logistic regression is one of the most popular Machine Learning algorithms, which comes under the Supervised Learning technique. It is used for predicting the categorical dependent variable using a given set of independent variables.
- Logistic regression predicts the output of a categorical dependent variable. Therefore the outcome must be a categorical or discrete value. It can be either Yes or No, 0 or 1, true or False, etc. but instead of giving the exact value as 0 and 1, **it gives the probabilistic values which lie between 0 and 1.**
- Logistic Regression is much similar to the Linear Regression except that how they are used. Linear Regression is used for solving Regression problems, whereas **Logistic regression is used for solving the classification problems.**
- In Logistic regression, instead of fitting a regression line, we fit an "S" shaped logistic function, which predicts two maximum values (0 or 1).
- The curve from the logistic function indicates the likelihood of something such as whether the cells are cancerous or not, a mouse is obese or not based on its weight, etc.
- Logistic Regression is a significant machine learning algorithm because it has the ability to provide probabilities and classify new data using continuous and discrete datasets.
- Logistic Regression can be used to classify the observations using different types of data and can easily determine the most effective variables used for the classification.
- The sigmoid function is a mathematical function used to map the predicted values to probabilities.
- It maps any real value into another value within a range of 0 and 1.

The below image is showing the logistic function:



Support Vector Machine

Support Vector Machine or SVM is one of the most popular Supervised Learning algorithms, which is used for Classification as well as Regression problems. However, primarily, it is used for Classification problems in Machine Learning. The goal of the SVM algorithm is to create the best line or decision boundary that can segregate n-dimensional space into classes so that we can easily put the new data point in the correct category in the future. This best decision boundary is called a hyperplane. SVM chooses the extreme points/vectors that help in creating the hyperplane. These extreme cases are called as support vectors, and hence algorithm is termed as Support Vector Machine. The data points or vectors that are the closest to the hyperplane and which affect the position of the hyperplane are termed as Support Vector. Since these vectors support the hyperplane, hence called a Support vector. Consider the below diagram in which there are two different categories that are classified using a decision boundary or hyperplane:



K-Nearest Neighbor

- K-Nearest Neighbour is one of the simplest Machine Learning algorithms based on Supervised Learning technique.
- K-NN algorithm assumes the similarity between the new case/data and available cases and put the new case into the category that is most similar to the available categories.
- K-NN algorithm stores all the available data and classifies a new data point based on the similarity. This means when new data appears then it can be easily classified into a well suite category by using K- NN algorithm.
- K-NN algorithm can be used for Regression as well as for Classification but mostly it is used for the Classification problems.
- K-NN is a **non-parametric algorithm**, which means it does not make any assumption on underlying data.

- It is also called a **lazy learner algorithm** because it does not learn from the training set immediately instead it stores the dataset and at the time of classification, it performs an action on the dataset.
- KNN algorithm at the training phase just stores the dataset and when it gets new data, then it classifies that data into a category that is much similar to the new data.
- **Example:** Suppose, we have an image of a creature that looks similar to cat and dog, but we want to know either it is a cat or dog. So for this identification, we can use the KNN algorithm, as it works on a similarity measure. Our KNN model will find the similar features of the new data set to the cats and dogs images and based on the most similar features it will put it in either cat or dog category.

User ID	Gender	Age	EstimatedSalary	Purchased
15624510	Male	19	19000	0
15810944	Male	35	20000	0
15668575	Female	26	43000	0
15603246	Female	27	57000	0
15804002	Male	19	76000	0
15728773	Male	27	58000	0
15598044	Female	27	84000	0
15694829	Female	32	150000	1
15600575	Male	25	33000	0
15727311	Female	35	65000	0
15570769	Female	26	80000	0
15606274	Female	26	52000	0
15746139	Male	20	86000	0
15704987	Male	32	18000	0
15628972	Male	18	82000	0
15697686	Male	29	80000	0
15733883	Male	47	25000	1
15617482	Male	45	26000	1
15704583	Male	46	28000	1
15621083	Female	48	29000	1
15649487	Male	45	22000	1
15736760	Female	47	49000	1

CONCLUSION AND FUTURE SCOPE

The handwritten digit recognition using convolutional neural network has proved to be of a fairly good efficiency. It works better than any other algorithm, including Artificial Neural Networks. In this project, we have successfully built a Python machine learning project on handwritten digit recognition app. We have built and trained the Convolutional neural network which is very effective for image classification purposes. Later on, we build the GUI where we draw a digit on the canvas then we classify the digit and show the results.

The proposed recognition system is implemented on handwritten digits taken from MNIST database. Handwritten digit recognition system can be extended to a recognition system that can also able to recognize handwritten character and handwritten symbols. Future studies might consider on hardware implementation of recognition system. This project is a very much preliminary project based on those .This world entitles the work of Google everyday who himself hasn't achieved that much data also. This project entitles some different new ideas on 1. Image Processing 2. Machine learning 3. Activation Functions 4. Statistical predictive modeling 5. Optimiser into the programing 6. Text analysis 7. Digit extraction Features. The future development of the applications based on algorithms of deep and machine learning is practically boundless.

Department of Electronics and Communication Engineering

Vision

To be recognized by the society at large as a full- fledged department, offering quality higher education in the Electronics and Communication Engineering field with research focus catering to the needs of the stakeholders and staying in tune with the advancing technological revolution and cultural changes.

Mission

To achieve the vision, the department will

- Establish a unique learning environment to enable the students to face the challenges in Electronics and Communication Engineering field.
- Promote the establishment of centres of excellence in niche technology areas to nurture the spirit of innovation and creativity among faculty and students.
- Provide ethical and value-based education by promoting activities addressing the societal needs.
- Enable students to develop skills to solve complex technological problems and provide a framework for promoting collaborative and multidisciplinary activities.

