

Preprocessing and Segmentation of Retinal Blood Vessels in Fundus Images using U-Net

Mrs. R Sudha Abirami¹ and Dr. G Suresh Kumar²

¹Research Scholar, ²Assistant Professor

Department of Computer Science and Engineering, Pondicherry University Karaikal Campus, Karaikal,
Email: rsudhaabirami27@gmail.com, mgsureshkumar@gmail.com

Abstract— Deep Learning plays an important role today in disease detection and prediction. All deep learning models need to be trained to process the input; Extract features and return prediction results. Before classification and prediction, the given input must be preprocessed to perform segmentation using augmentation. Only with the help of preprocessed images each model can make accurate predictions at higher speeds. This proposed work aimed to detect Diabetic Eye Diseases by means of segmenting the augmented images using U-Net. U-Net is familiar with its Encoder-Decoder architecture for sampling. Retinal Blood Vessel is one of the most precise parts of an eye. Based on the nature of this blood vessel one can identify whether it is affected by diabetic retinopathy or not. So, segmenting the blood vessel helps to classify the disease category in early stage and of course U-Net is probably meant for segmenting medical images. In this paper, the discussions will be made on preprocessing eye images from the data set, segmenting those images using U-Net to extract the retinal blood vessel, classification based on segmentation..

Index Terms— Deep Learning, Prediction, Classification, Segmentation, Augmentation, U-Net, Diabetic Eye Diseases.

I. INTRODUCTION

A. Image Processing:

Image processing as the name suggests, means processing images, and many techniques are required to reach the goal. It is the core domain of computer vision that plays a key role in many real-world examples, such as robotics, self-driving cars, and object recognition. Image processing allows us to transform and manipulate thousands of images simultaneously and derive useful insights from them. It has a wide range of applications in almost every field. The final output may be in the form of either an image or an equivalent feature of that image. This can be used for further analysis and decision making. An image can be represented as a 2D function $F(x, y)$ where x and y are the spatial coordinates. The amplitude of F at a particular value of x, y is known as the intensity of the image at that point. If the x, y and amplitude values are finite, it can be called as digital image. This is an array of pixels arranged in columns and rows. A pixel is an image element that contains information about intensity and color. Images can also be represented in 3D, where x, y , and z are spatial coordinates. Pixels are arranged in a matrix. This is called an RGB image. There are different types of images: i) RGB Image - Contains three layers of a 2D image, these layers being the red, green and blue channels, ii) Grayscale Images - These images contain shades of black and white and contain only one channel.

B. Segmentation:

Image segmentation is a computer vision task that divides an image into regions by assigning a label to each pixel in the image. It provides more information about the image than object detection, which draws bounding boxes around detected objects, or image classification, which assigns labels to objects. Segmentation is useful and can be used in real-world applications such as medical imaging, clothing segmentation, flood maps, and self-driving cars. There are two types of image segmentation.

- Semantic Segmentation: Classify each pixel with a label.
- Instance Segmentation: Classifies each pixel to distinguish each object instance.

In Ref. [1], U-Net is a semantic segmentation method originally proposed for medical image segmentation. This was one of the early deep learning segmentation models, and the U-Net architecture (Figure 1) is also used in many of his GAN variants such as his Pix2Pix generator.

In Ref. [2], the U-Net is an elegant architecture that solves most of the occurring issues. It uses the concept of fully convolutional networks for this approach. The intent of the U-Net is to capture both the features of the context as well as the localization. This process is completed successfully by the type of architecture built. The main idea of the implementation is to utilize successive contracting layers, which are immediately followed by the up-sampling operators for achieving higher resolution outputs on the input images.

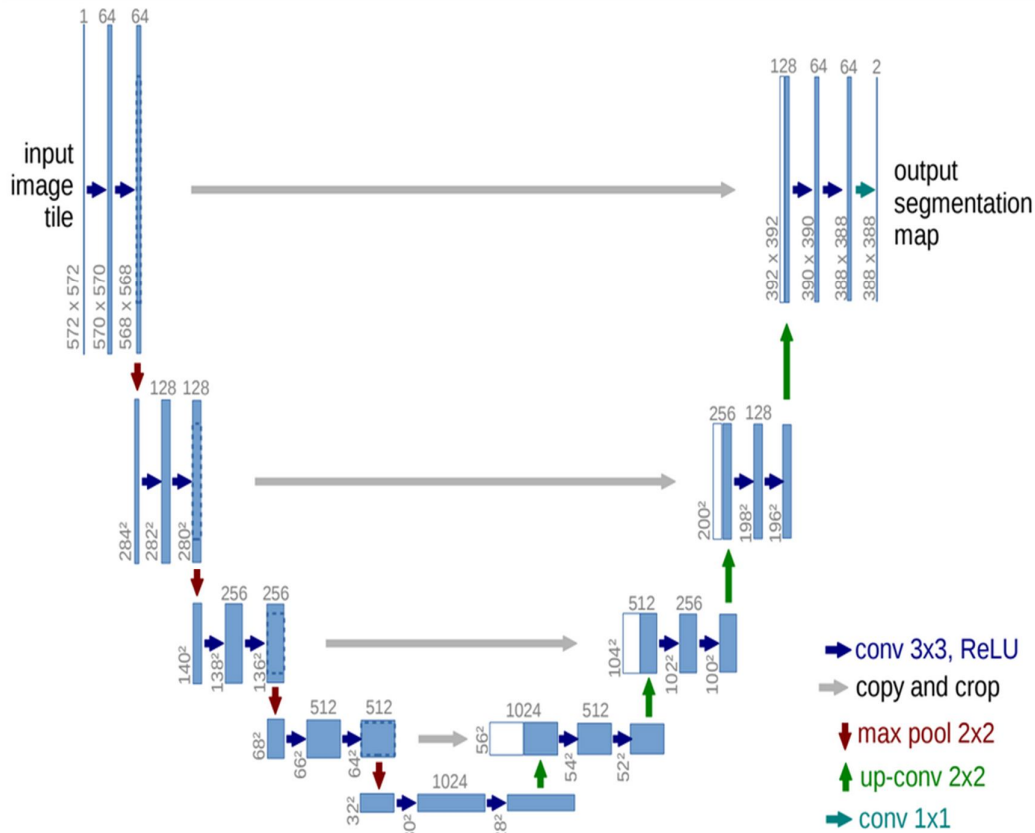


Figure1.U-Net Architecture

This paper is arranged as chapters in the following manner, Chapter II contains Literature Survey, Chapter III contains Image Preprocessing, Chapter IV contains Segmentation, Chapter V contains Proposed Method, Chapter VI contains Segmentation result and discussions, and Chapter VII contains Conclusion.

II. LITERATURE SURVEY

In Ref. [4], a typical CNN has a multi-layered structure such as feed forward networks. As different from feed forward networks, a CNN can include several Convolutional layers with a sub-sampling section. All the

parameters have been calculated for the CNN after three image processing stages as resizing the image, applying Histogram Equalization, and applying CLAHE. After the image processing-based enhancement, the classification was made using the CNN. The performance of the introduced method was evaluated by using 400 retinal fundus images in the MESSIDOR database. In Ref. [5], this work depicts the green channel of the RGB model exhibits the best contrast between the vessels and background while the red and blue ones tend to be noisier. The grey image from the green channel is processed and the retinal blood vessels appear darker in the grey image and then invert it to appear brighter than non-vessel background. Salt and pepper noise is added in order to represent the presence of noise. In order to remove the salt and pepper noise, order and median filters are used. The output of the order filter gives better contrast between the vessels and the background, thereby removing the noise more accurately than the other filters.

In Ref. [6], This study aimed to detect Optic Disc. Some of the features of Diabetic Retinopathy are exudates, hemorrhages and micro aneurysms. Detection and removal of optic disc plays a vital role in extraction of these features. This paper focuses on detection of optic disc using various image processing techniques, algorithms such as canny edge, Circular Hough (CHT). Retinal images from IDRiD, Diaret_db0, Diaret_db1, Chasedb and Messidor datasets were used.

In Ref. [7], the proposed model has been trained with three types, back propagation NN, Deep Neural Network (DNN) and Convolutional Neural Network (CNN) after testing models with CPU trained Neural network gives lowest accuracy because of one hidden layer whereas the deep learning models are outperforming NN. The Deep Learning models are capable of quantifying the features as blood vessels, fluid drip, exudates, hemorrhages and micro aneurysms into different classes. In Ref. [8], Glaucoma is a group of conditions, in which high pressure inside the eye damages the optic nerve of the eye. The vision lost due to glaucoma is irreversible and cannot be regained. Hence it is very important to detect this disease as early as possible and treat early to preserve vision. In this paper, the performance of five preprocessing techniques is compared namely Contrast Adjustment, Adaptive Histogram equalization, Median filtering, Average filtering and Homomorphic filtering. The performances of these techniques are evaluated using Mean Square Error (MSE) and Peak Signal to Noise Ratio (PSNR).

In Ref. [9], proposed automated methods consist of pre-processing, blood vessels extraction, optic disc segmentation and macula region segmentation. Initially, pre-processing is performed using shade correction and top-hat transformation for enhancement of dark anatomical structures such as blood vessels and macula/fovea region. A novel graph cut method is used to extract blood vessels. Then template based matching and morphological operations are used for detection and extraction of optic disc. Finally, post processing is used for detection of macula in retinal images. In Ref. [10], the work proposed was an OD segmentation model from fundus images based on Retina Net extension with DenseNet that addresses the vanishing gradient problem, enhances feature propagation, performs deep supervision, strengthens feature reuse and reduces the number of parameters. The model was developed based on promising results achieved by the Retina Net and the DenseNet in many object detection problems. Combining both models facilitates the reuse of computation through dense connections and improves gradient flow.

In Ref. [11], segmenting the optic disc (OD) is an important and essential step in creating a frame of reference for diagnosing optic nerve head pathologies such as glaucoma. The main contribution of this paper is in presenting a novel OD segmentation algorithm based on applying a level set method on a localized OD image. To prevent the blood vessels from interfering with the level set process, an in-painting technique was applied. The new automatic eye disease diagnosis system has to be robust, fast, and highly accurate, in order to support high workloads and near-real-time operation.

III. IMAGE PREPROCESSING

The purpose of pre-processing is to raise the image's quality so that we can analyze it more effectively. Preprocessing allows us to eliminate unwanted distortions and improve specific qualities that are essential for the application we are working on. Those characteristics could change depending on the application. There are four different types of Image Pre-Processing techniques and they are as follows;

1. Pixel brightness transformations/ Brightness corrections
2. Geometric Transformations
3. Image Filtering and Segmentation
4. Fourier transform and Image re-saturation

The brightness of a pixel is altered by *Brightness transformations*, which are dependent on the characteristics of the individual pixel. In PBT, the value of the output pixel depends only on the value of the matching input pixel.

Enhancing contrast is a crucial component of image processing for both human and machine vision. It is commonly used in speech recognition, texture synthesis, medical image processing, and many other image/video processing applications as a pre-processing step.

The most common Pixel brightness transforms operations are

1. Gamma correction or Power Law Transform
2. Sigmoid stretching
3. Histogram equalization

Two commonly used point processes are multiplication and addition with a constant.

$$g(x) = \alpha f(x) + \beta \quad (1)$$

The parameters $\alpha > 0$ and β are called the gain and bias parameters and sometimes these parameters are said to control contrast and brightness respectively.

A. Histogram Equalization

It is a well-known contrast enhancement technique due to its performance on almost all types of images. Histogram equalization provides a sophisticated method for modifying the dynamic range and contrast of an image by altering that image such that its intensity histogram has the desired shape. Unlike contrast stretching, histogram modeling operators may employ non-linear and non-monotonic transfer functions to map between pixel intensity values in the input and output images.

The normalized histogram can be represented as,

$$P(n) = (\text{number of pixels with intensity } n) / (\text{total number of pixels})$$

B. Image Filtering and Segmentation

The purpose of utilizing filters is to change or improve the qualities of the images and/or to extract important data from the images, such as edges, corners, and blobs. A kernel, which is a tiny array applied to each pixel and its neighbors inside a picture, defines a filter. Some of the basic filtering techniques are; i. Low Pass Filtering (Smoothing), ii. High pass filters (Edge Detection, Sharpening), iii. Directional Filtering, iv. Laplacian Filtering. Often based on the properties of the picture's pixels, *Image Segmentation* is a widely used method in digital image processing and analysis to divide an image into various parts or areas. Foreground and background can be distinguished in an image using segmentation, and pixels can be grouped together according to their similarity in color or shape. Image Segmentation is mainly used in Face detection, medical imaging, Machine vision, Autonomous Driving.

C. Fourier Transform

In Ref. [3] is an important image processing tool used to decompose an image into its sine and cosine components. The output of the transform represents an image in the Fourier or frequency domain, while the input image is equivalent in the spatial domain. In Fourier-domain images, each point represents a specific frequency in the spatial-domain image. Fourier transforms are used in a variety of applications such as image analysis, image filtering, image reconstruction, and image compression.

IV. IMAGE SEGMENTATION

A. What is Segmentation?

Segmentation is the separation of one or more regions or objects in an image based on a discontinuity or a similarity criterion. In Ref. [21], A region in an image can be defined by its border (edge) or its interior, and the two representations are equal.

Segmentation approaches:

- *Pixel-based segmentation:* each pixel is segmented based on gray-level values, no contextual information, only histogram. Example: Thresholding.
- *Region-based segmentation:* considers gray-levels from neighboring pixels by – including similar neighboring pixels (region growing), – split-and-merge, – or watershed segmentation.
- *Edge-based segmentation:* Detects and links edge pixels to form contours.

Following are the primary types of image segmentation techniques: (in Figure 2)

1. Edge-Based Segmentation
2. Thresholding Segmentation
3. Region-Based Segmentation
4. Clustering-Based Segmentation

5. Watershed Segmentation

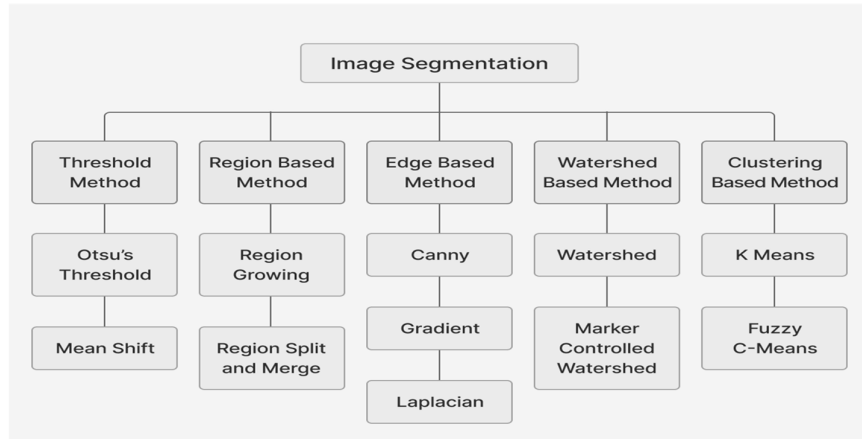


Figure 2. Segmentation approaches

1. **Edge-Based Segmentation:** Edge-based segmentation is a popular image processing technique that identifies the edges of various objects in a given image. It helps locate features of associated objects in the image using the information from the edges. Edge detection helps strip images of redundant information, reducing their size and facilitating analysis.

Edge-based segmentation algorithms identify edges based on contrast, texture, color, and saturation variations. They can accurately represent the borders of objects in an image using edge chains comprising the individual edges.

2. **Thresholding Segmentation:** Thresholding is the simplest image segmentation method, dividing pixels based on their intensity relative to a given value or threshold. It is suitable for segmenting objects with higher intensity than other objects or backgrounds.

The threshold value T can work as a constant in low-noise images. In some cases, it is possible to use dynamic thresholds. Thresholding divides a grayscale image into two segments based on their relationship to T , producing a binary image.

3. **Region-Based Segmentation:** Region-based segmentation involves dividing an image into regions with similar characteristics. Each region is a group of pixels, which the algorithm locates via a seed point. Once the algorithm finds the seed points, it can grow regions by adding more pixels or shrinking and merging them with other points.

4. **Clustering-Based Segmentation:** Clustering algorithms are unsupervised classification algorithms that help identify hidden information in images. They augment human vision by isolating clusters, shadings, and structures. The algorithm divides images into clusters of pixels with similar characteristics, separating data elements and grouping similar elements into clusters.

5. **Watershed Segmentation:** Watersheds are transformations in a grayscale image. Watershed segmentation algorithms treat images like topographic maps, with pixel brightness determining elevation (height). This technique detects lines forming ridges and basins, marking the areas between the watershed lines. It divides images into multiple regions based on pixel height, grouping pixels with the same gray value. The watershed technique has several important use cases, including medical image processing. For example, it can help identify differences between lighter and darker regions in an MRI scan, potentially assisting with diagnosis.

B.CNN for Image Segmentation:

In Ref. [12], we applied Convolutional Neural Networks (CNN) on the semantic segmentation of remote sensing images. As well as, improving the Encoder- Decoder CNN structure SegNet with index pooling and U-Net to make them suitable for multi-targets semantic segmentation of remote sensing images.

Convolutional neural network is a hierarchical model whose input is raw data, such as RGB image and raw audio data. Convolutional neural networks extract high level semantic information layer by layer from the input layer of raw data by stacking a series of operations such as convolution operation, pooling operation and mapping of non-linear activation functions. This process is called “feed-forward operation”. Different types of operations in the convolutional neural networks are called “layers”. Convolution operations are convolutional layers and

pooling operations are pooling layers. The last layer of convolutional neural network transforms its target tasks (classification, regression, etc.) into the objective function. By calculating the error or loss between the predicted value and the real value, the error or loss back-forward layer by layer by the back-propagation algorithm to update the parameters of every layer and then back-forward again and again until the network model converges. In Ref. [13], during past days, a lot of research has been carried out for retinal blood vessel segmentation for identification of Diabetic Retinopathy using various machine learning and deep learning models. In this research work, Convolutional Neural Network (CNN) and CLAHE are applied together to tackle the problem of retinal blood vessel segmentation of images over the DRIVE dataset. The method undergoes pre-processing- grey scale conversion and CLAHE, feature extraction using morphological feature, segmentation, training and prediction using CNN. Experimental evaluation shows that the proposed method outperforms 98.06% accuracy.

In Ref. [14], a comprehensive review of the literature has been written, covering a broad spectrum of pioneering works for semantic and instance-level segmentation, including fully convolutional pixel-labeling networks, encoder-decoder architectures, multi-scale and pyramid-based approaches, recurrent networks, visual attention models, and generative models in adversarial settings. We investigate the similarity, strengths and challenges of these deep learning models, examine the most widely used datasets, report performances, and discuss promising future research directions in this area. In Ref. [15], with the advent of neural networks, deep convolutional neural networks (DCNNs) provide benchmarking results in the problems related to computer vision. Manifold DCNNs have been proposed for semantic segmentation such as U-Net, DeepU-Net, ResUNet, DenseNet, RefineNet, etc. The general procedure is common for all the models. It has three phases - pre-processing, processing and output generation. The outputs of the processing phase are the masked image and segmented image. In this paper, a systematic critique of the existing DCNNs for semantic segmentation has been manifested. The datasets and the architectures of the existing models have also been discussed in this paper with illustrations.

In Ref. [16], the proposed model has 13 layers and uses dilated convolution and max-pooling to extract small features. Ghost model deletes the duplicated features, makes the process easier, and reduces the complexity. The Convolution Neural Network (CNN) generates a feature vector map and improves the accuracy of area or bounding box proposals. Restructuring is required for healing. As a result, convolutional neural networks segment medical images. It is possible to acquire the beginning region of a segmented medical image. The proposed model gives better results as compared to the traditional models, it gives an accuracy of 96.05, Precision 98.2, and recall 95.78. In Ref. [17], they present a network and training strategy, that relies on the strong use of data augmentation to use the available annotated samples more efficiently. The architecture consists of a contracting path to capture context and a symmetric expanding path that enables precise localization. We show that such a network can be trained end-to-end from very few images and outperforms the prior best method (a sliding-window convolutional network) on the ISBI challenge for segmentation of neuronal structures in electron microscopic stacks.

In Ref. [18], after researching various techniques, they have found that, the CNN is one the most powerful tool in image segmentation technique. Detailed analysis of CNN is also done here explaining different layers and workings of each layer. As we know CNN technology is at a boost of implementation nowadays in making the human life more and more convenient and less manual. There have already been a lot of work done in various fields like commutation, medical tasks, crop monitoring, road transportation, activity detection, product quality monitoring. In Ref. [19], a novel attention Gabor network (AGNet) based on deep learning for medical image segmentation that is capable of automatically paying more attention to the edge and consistently for improvement to the segmentation performance is proposed. The proposed model consists of two components. The first one is to determine the approximate location of the organs of interest in the image using convolution filters, and the other one is to highlight salient edge features intended for a specific segmentation task using Gabor filters. In order to facilitate collaboration in between the two parts, a region attention mechanism based on Gabor maps is suggested. The mechanism improved performance by learning to focus on the salient regions of the image that are useful for the authors' tasks.

In Ref. [20], fully automatic segmentation of wound areas in natural images is an important part of the diagnosis and care protocol since it is crucial to measure the area of the wound and provide quantitative parameters in the treatment. Various deep learning models have gained success in image analysis including semantic segmentation. This manuscript proposes a novel convolutional framework based on MobileNetV2 and connected component labeling to segment wound regions from natural images. The advantage of this model is its lightweight and less compute-intensive architecture. The performance is not compromised and is comparable to deeper neural networks. We build an annotated wound image dataset consisting of 1109-foot ulcer images from 889 patients to train and test the deep learning models.

V. PROPOSED METHOD

A. Dataset:

In the proposed method, the UNET model is trained with DRIVE dataset. This dataset contains 20 training images with mask and manual ground truth images. The images were enhanced using some preprocessing techniques such as, Grayscale transformations, Brightness transformations, applying CLAHE (Contrast Limited Adaptive Histogram Equalization), Extracting edges by Canny Edge Detection etc., The aim of this method is to segment the blood vessel from retina for classification of disease severity.

B. Flow Diagram:

The flow diagram (in Figure 3) depicts the work flow of process to extract blood vessel and segment them using UNET Model. In the proposed method, images are acquired from the dataset. The RGB image is then converted into a grayscale image. The Gray image is fed to the process of CLAHE (Contrast Limited Adaptive Histogram Equalization). The image is processed with Canny Edge Detection method. By this, Retinal Vessels are extracted from the image. The output is considered as ground truth image for segmentation process. The UNET model is trained to segment the retinal blood vessel from the given image. This model is implemented in Google Colab, NVIDIA Tesla K80 GPU. The model performance is measured using segmented image by comparing them with ground truth.

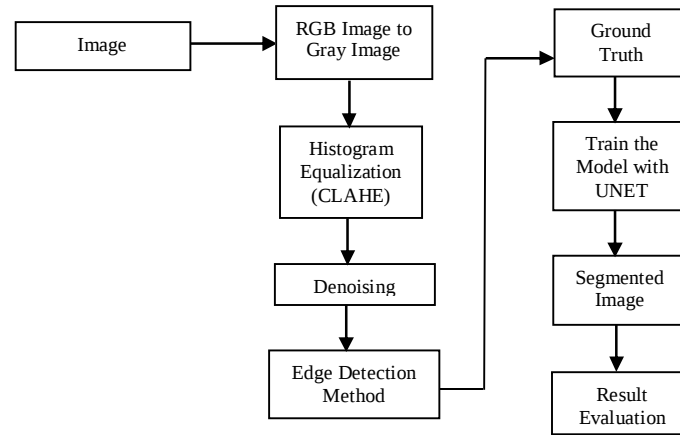


Figure 3. Flow Diagram of Proposed Method

C. ALGORITHM:

Input: I_i , Read images from Dataset,

Process:

Step1: Convert RGB Image to Grayscale Image

Step2: Increase Brightness using CLAHE.

Step3: Denoise the image by Non-Local Mean Denoising function.

Step4: Detect the edges by Canny Edge Detector – Ground Truth Image, GI .

Step5: Train the model, $U_i = I_i + GI_i$.

Step6: Segmented Image, S_i .

Output: Retinal Blood Vessels Extraction.

D. CLAHE (Contrast Limited Adaptive Histogram Equalization):

The CLAHE operates on small regions in the image, called tiles, rather than the entire image. The neighboring tiles are then combined using bilinear interpolation to remove the artificial boundaries. This algorithm can be applied to improve the contrast of images. There are two parameters to be considered. They are,

- clip Limit – This parameter sets the threshold for contrast limiting. The default value is 40.
- tileGridSize – This sets the number of tiles in the row and column. By default, this is 8×8 . It is used while the image is divided into tiles for applying CLAHE.

E. De-noising Images:

In Ref. [22], its basic idea is to build a pointwise estimation of the image, where each pixel is obtained as a weighted average of pixels centered at regions that are similar to the region centered at the estimated pixel. For a given pixel x_i in an image x , NLM (x_i) indicates the NLM-filtered value.

Let $w_{i,j}$ be the weight of x_j to x_i , which is computed by.

$$W_{i,j} = 1/C_i \exp (- \|X_i - X_j\|_2^2) / h \quad (2)$$

where C_i denotes a normalization factor, and h indicates a filter parameter.

F. Canny Edge Detection Method:

This method contains four major steps to process;

- Reduce Noise using Gaussian Smoothing.
- Compute image gradient using Sobel filter.
- Apply Non-Max Suppression or NMS to just keep the local maxima
- Finally, apply Hysteresis thresholding which that 2-threshold values T_{upper} and T_{lower} which is used in the Canny () function.

Figure 4. shows the original image, preprocessed, edge image, and segmented image.

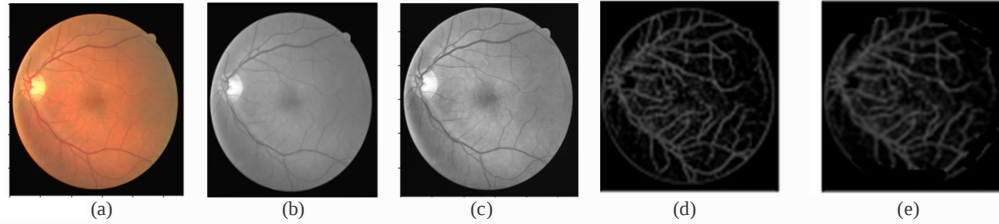


Figure 4. (a) Original image, (b) Grayscale image, (c) CLAHE image, (d) Canny Edge image and (e) Segmented image.

VI. SEGMENTATION RESULTS

After preprocessing, the U-Net model is trained with images and masks. Original images and Ground truth images are given as training and validation inputs. Images are given to x coordinate and masks are given to y coordinate to perform 2D convolutions. The DRIVE dataset contains manually annotated masks considered as ground truth images. The model performs out well in all epochs. The number of epochs started from 5, and ends with 100. The following table (Table.1) and Figure 5. shows the accuracy achieved for each number of epochs.

TABLE.1. NO. OF EPOCHS VS ACCURACY

No. of Epochs	Accuracy (%)	No. of Epochs	Accuracy (%)
	88.73	50	91.23
10	88.74	70	94.71
20	89.31	100	98.92

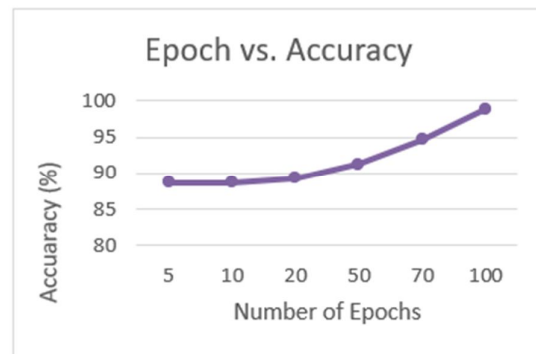


Fig.5 No. of Epochs and Accuracy

A. Result Comparison:

In previous research works, the authors used Convolutional Neural Network models and H-minima for classification using the DRIVE dataset. But segmentation helps the classification task in better way. And also, U-Net models is specifically meant for bio-medical images. This model really predicts the segmented images as in

the ground truth image. The edge images were generated automatically. This method is implemented in Google Colab with GPU environment. From the results, the proposed model is compared with existing works of different datasets and different models. The proposed method outperforms well with U-Net Model. The following table (Table.2) shows the comparison about the dataset used and accuracy obtained by other authors.

TABLE II COMPARISON OF PROPOSED METHOD WITH PREVIOUS WORK ON DRIVE DATASET

Author	Dataset	Model	Accuracy (%)
[13]	DRIVE	CNN	98.06
[23]	DRIVE	CNN+UNET	97.90
[24]	STARE & DRIVE	H-Minima	95.91 & 96.72
Proposed Method	DRIVE	CANNY Edge Detection + UNET	98.92

VII. CONCLUSION

In the proposed method, DRIVE dataset is used and Canny Edge Detection with U-Net model was implemented. The model uniquely identifies all the edges present in the original eye image, also it represents the blood vessels of retina. From the identified edges, the U-Net model is trained with images and masks for various number of epochs. Finally, it achieves better accuracy compared to other previous works. This work can be enhanced furtherly as, by applying different datasets such as, IDRiD, Messidor-2, STARE etc., and the segmentation can also be done with other models like Attention U-Net, Res U-Net. Due to its easy-to-access and good performance, the proposed method accelerates the diagnosis of Diabetic Eye Diseases in early stages itself.

REFERENCES

- [1] <https://pyimagesearch.com/2022/02/21/u-net-image-segmentation-in-keras/>
- [2] <https://blog.paperspace.com/unet-architecture-image-segmentation/>
- [3] <https://www.mygreatlearning.com/blog/introduction-to-image-pre-processing/>
- [4] D. Jude Hemanth¹, Omer Deperlioglu², Utku Kose³, “An enhanced diabetic retinopathy detection and classification approach using deep convolutional neural network”, <https://doi.org/10.1007/s00521-018-03974-0>.
- [5] P. Pearlina Sheeba, V. Radhamani, “Analyzing and Feature Extraction of Diabetic Retinopathy in Retinal Images”, International Journal of Engineering Research & Technology (IJERT) <http://www.ijert.org>, ISSN: 2278-0181.
- [6] Priyanka Konatham, Mounika Venigalla, Lakshmi Pooja Amaraneni, K. Suvarna Vani, “Automatic Detection of Optic Disc for Diabetic Retinopathy”, International Journal of Innovative Technology and Exploring Engineering (IJITEE), ISSN: 2278-3075 (Online), Volume-9 Issue-7, May 2020.
- [7] Suvajit Dutta, Bonthala CS Manideep, Syed Muzamil Basha, Ronnie D. Caytiles¹ and N. Ch. S. N. Iyengar, “Classification of Diabetic Retinopathy Images by Using Deep Learning Models”, <http://dx.doi.org/10.14257/ijgdc.2018.11.1.09>.
- [8] Mrs.S.Rathinam*¹ M.E.(Ph.D), Dr.S.Selvarajan*² M.E.,Ph.D., “Comparison of Image Preprocessing Techniques on Fundus Images for Early Diagnosis of Glaucoma”, International Journal of Scientific & Engineering Research, Volume 4, Issue 12, December-2013 1368 ISSN 2229-5518.
- [9] P. R. Wankhede, K. B. Khanchandani, “Feature Extraction In Retinal Images Using Automated Methods”, International Journal Of Scientific & Technology Research Volume 9, Issue 03, March 2020 ISSN 2277-8616.
- [10] Manal AlGhamdi, “Optic Disc Segmentation in Fundus Images with Deep Learning Object Detector”, Journal of Computer Science 2020, 16 (5): 591.600 DOI: 10.3844/jcssp.2020.591.600.
- [11] Ahmed Almazroa, Weiwei Sun, Sami Alodhayb, Kaamran Raahemifar, Vasudevan Lakshminarayanan, “Optic disc segmentation for glaucoma screening system using fundus images”, <https://doi.org/10.2147/OPHTH.S140061>.
- [12] Muhammad Alam, Jian-Feng Wang, Cong Guangpei, LV Yunrong, Yuanfang Chen, “Convolutional Neural Network for the Semantic Segmentation of Remote Sensing Images”, Mobile Networks and Applications (2021) 26:200–215, <https://doi.org/10.1007/s11036-020-01703-3>.
- [13] Arun Kumar Yadav, Arti Jain, Jorge Luis Morato Lara and Divakar Yadav, “Retinal Blood Vessel Segmentation using Convolutional Neural Networks”, DOI: 10.5220/0010719500003064.
- [14] Shervin Minaee, Yuri Boykov, Fatih Porikli, Antonio Plaza, Nasser Kehtarnavaz, and Demetri Terzopoulos, “Image Segmentation Using Deep Learning: A Survey”, arXiv: 2001.05566v5 [cs.CV] 15 Nov 2020.
- [15] Rishipal Singh, Rajneesh Rani, “Semantic Segmentation using Deep Convolutional Neural Network: A Review”, International Conference On Intelligent Communication And Computational Research (ICICCR-2020), <https://ssrn.com/abstract=3565919>.
- [16] Marcelo Zambrano-Vizuete, Miguel Botto-Tobar, Carmen Huerta-Sua´rez, Wladimir Paredes-Parada, Darwin Patiño Perez, Tariq Ahamed Ahanger, and Neilsy Gonzalez, “Segmentation of Medical Image Using Novel Dilated Ghost Deep Learning Model”, Computational Intelligence and Neuroscience, <https://doi.org/10.1155/2022/6872045>.

- [17] Olaf Ronneberger, Philipp Fischer, and Thomas Brox, "U-Net: Convolutional Networks for Biomedical Image Segmentation", arXiv: 1505.04597v1 [cs.CV] 18 May 2015.
- [18] Ravi Kaushik, Shailender Kumar, "Image Segmentation Using Convolutional Neural Network", International Journal Of Scientific & Technology Research Volume 8, Issue 11, November 2019, ISSN 2277-8616.
- [19] Shaoqiong Huang, Mengxing Huang, Yu Zhang, Jing Chen, Uzair Bhatti, "Medical image segmentation using deep learning with feature enhancement", IET Image Processing, doi: 10.1049/iet-ipr.2019.0772.
- [20] Chuanbo Wang, D. M. Anisuzzaman, Victor Williamson, Mrinal Kanti Dhar, Behrouz Rostami, Jeffrey Niezgoda, Sandeep Gopalakrishnan, Zeyun Yu, "Fully automatic wound segmentation with deep convolutional neural networks", <https://doi.org/10.1038/s41598-020-78799-w>.
- [21] <https://datagen.tech/guides/image-annotation/image-segmentation/>, <https://www.v7labs.com/blog/image-segmentation-guide>.
- [22] Linwei Fan, Fan Zhang, Hui Fan and Caiming Zhang, "Brief review of image denoising techniques", Visual Computing for Industry, Biomedicine, and Art (2019) 2:7 <https://doi.org/10.1186/s42492-019-0016-7>.
- [23] Wang Xiancheng, Li Weia, Miao Bingyi, Jing Hed, Zhangwei Jiangf, Wen Xue, Zhenyan Ji g, Gu Hongh, Shen Zhaomeng, "Retina Blood Vessel Segmentation Using A U-Net Based Convolutional Neural Network", International Conference on Data Science (ICDS 2018), www.sciencedirect.com.
- [24] Boubakar Khalifa Albargathe, S.M., Kamberli, E., Kandemirli, F. et al. Blood vessel segmentation and extraction using H-minima method based on image processing techniques. Multimedia Tools Appl 80,2565–2582(2021), <https://doi.org/10.1007/s11042-020-09646-3>.