

Advanced deep learning approaches for early detection and localization of ocular diseases

Ali Mohammed Ridha^{1*}, Mohammed Jamal Mohammed², Hussban Abood Saber³, Mustafa Habeeb

Chyad⁴, Maryam Hussein Abdulameer⁵

^{1,2}College of Medicine, University of Al-Ameed, 56001 Karbala, Iraq; eng.mohammed.j@alameed.edu.iq (A.M.R.)

alialmosawi@alameed.edu.iq (M.J.M.).

³Department of Electrical and Electronic Engineering, University of Kerbala, 56001 Karbala, Iraq;

hussban.a@s.uokerbala.edu.iq (H.A.S.).

^{4,5}University of Warith Al-Anbiyaa, 56001 Karbala, Iraq; mareim.hu@uowa.edu.iq (M.H.C.) mustfa1521995@gmail.com (M.H.A.).

Abstract: Rece with t advancements in modern technology have significantly enhanced the transmission of information, particularly in image processing, utilizing deep learning algorithms. This study aims to propose a a robust deep-learning strategy for detecting and recognizing eye defects and diseases from medical images. We present a detailed practical simulation of hybrid deep learning techniques designed for medical image classification based on multi-descriptor algorithms. The focus is on the classification of eye diseases by applying an advanced deep-learning algorithm to a dataset comprising various pathological eye conditions. Training operations for the proposed algorithm were conducted following the initialization phase, which included the extraction of multi-specification features. This enables the deep learning model to effectively analyze input eye images and accurately diagnose conditions. Our results demonstrate a diagnostic efficiency of 99%, with an error rate not exceeding 0.015%. The findings underscore the high efficiency and accuracy of deep learning algorithms in classifying and analyzing image data, thereby significantly reducing the workload for healthcare professionals.

Keywords: Deep learning techniques, Detect tire defects image classification mechanisms, Eye diseases, Fingerprint.

1. Introduction

Machine learning and profound learning are parts of artificial intelligence strategies that permit PC systems to advance straightforwardly from models, data, and experience. Numerous algorithms are accessible and sadly likewise unfit to determine the ideal algorithm for the issue. Every algorithm contrast in the idea of its design, its strategy for activity, and the kind of data it processes. This distinction is because of the multiplicity and variety of utilizations, the fluctuation of the data, and its huge number. Machines are not inherently astute. In any case, the machines, they are intended to perform unequivocal assignments, similar to action rail lines, traffic stream control, profound pit exhuming, go to space and shoot moving things. Machines play out their undertakings a ton faster and at a more elevated level exactness stood out from individuals. The fundamental differentiation among individuals and machines in taking consideration of their obligations is information. the human mind it gets data gathered by the five identifies: sight, hearing, smell, taste and contact. During the time spent discernment, data are coordinated, saw by contrasting them and past experiences that were put away in memory and unravelled. As necessary, the mind pursues the decision and guides the body parts to answer it this work. Close to the finish of the experience, it may be put away in memory for future benefits [1-5]. The innovations endeavour to address and re-enact the individual's information and human mental view of the individual's look so the impression is of an outside, three-layered nature. The intelligence of the individual human mind has been endeavoured for a really long time to perceive and

control vision caught through natural eyes within the human mind. These instincts act as the centre of growing new innovations. Rich assets presently accelerate crafted by examiners to find all the finer subtleties in the design of pictures [4,5]. Such upgrades are because of the abilities of representing the human mind through using of best systems and profound learning procedures, for example, neural cloud networks like CNN. Applications on Google, Facebook, Microsoft, as well as Snap/talk are delayed consequences of enormous enhancements in PC vision and profound learning. Over the long run, the making of a dream in light of mental discernment has transformed into a simple heuristic technique into insightful treatment systems that could explain this current reality [6]. Profound learning is described as a machine learning procedure, where many layers of data processing plants are used, through classification plans and their characteristics, or by learning by depiction. Profound learning is completed by neural networks, with many layers which is business as usual [7,8]. Nonetheless, it has become renowned lately, on account of three elements: First, expanded abilities to process (for instance video cards, designs processors, and so on); second, through computers; Third, in light of late advances and enhancements in profound learning research. Profound learning algorithms can be set up into three subgroups, which depend upon whether the algorithms are ready to convey the results Which may be gotten or not. The three subgroups are requested as unsupervised, administered, and hybrid [9, 10].

1.1. Features Extraction (Descriptor)

A feature descriptor is a method that extracts feature descriptions for a specific point of interest (or the whole image). Feature descriptors act as a kind of digital “fingerprint” that we can use to distinguish one feature from another by encoding information of interest into a string of numbers. A feature is a piece of data that is important to complete the calculations needed for a particular application. Features in an image can be specific elements such as points, edges, or objects. General scanning or feature detection applied to the image may also produce additional features that help in fully analyzing the image. Any point in n-dimensional space known as feature space contains a feature vector. For instance, a 2D feature space will have two inclination aspects. Then, we could put either sets of metrics on this two-dimensional space. Subsequent to giving each aspect an importance, we transform this locale into a portrayal space for all possible chains of implications. We characterize aspects as the quantity of data columns used to address attributes like location, color, density, neighborhood, and so forth. We should pick which aspects and columns to utilize while arranging or gathering data to obtain helpful data. The quantity of aspects in the feature space increments with the degree of detail required to address [11, 12].

1.2. Multi-Descriptor Features Extraction

The idea of a descriptor is generally to classify or describe a specific text or topic, especially when indexing or in an information retrieval system. In terms of computers, a descriptor is a data element that stores the attributes of some other predicate and is called a task descriptor. In computer vision, image descriptors describe elementary characteristics such as shape, color, texture, or movement of images. These are visual features of images that are expressed by dividing the image into small square pieces (pixels), and a number (weight) is assigned to each piece according to the image specifications, i.e., based on "On size, color, shade, resolution, depth and other image characteristics. Figure 1, illustrates a mapping diagram of the multi-Descriptor technique for image features extraction [10-15].

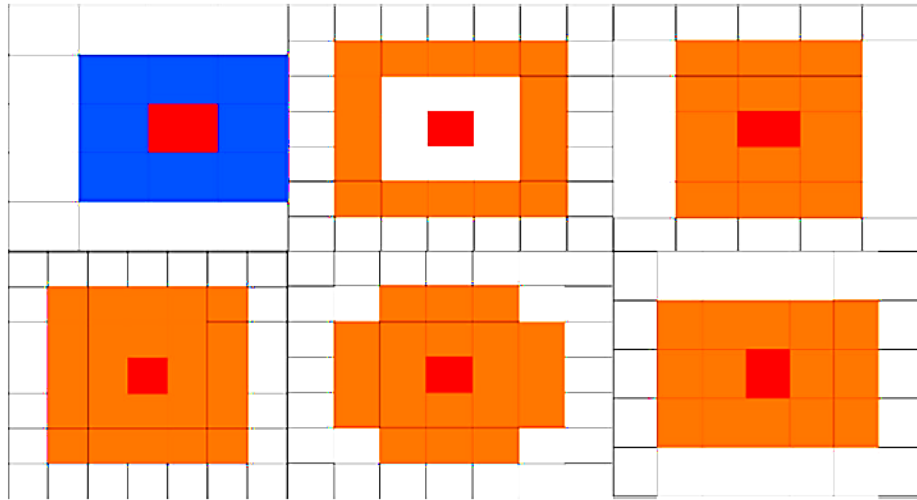


Figure 1.
Image features extraction mapping chart of the multi-Descriptor technique [10-15].

In computer vision, visual descriptors or image descriptors are descriptions of visual features of contents in images or videos, or algorithms or applications that produce such descriptions. They describe elementary properties such as shape, color, texture, or motion, among others.

1.3. Deep Convolutional Neural Networks (CNNs)

The CNN is a deep learning algorithm utilized for object recognition and features detection. The idea behind CNN is to set up the system to mimic human learning. The convolutional neural network algorithm consists of a group of interconnected layers that work to simulate the brain's neural network as much as possible. A huge dataset is used to train a CNN based on this correlation. The examples in this dataset are not only the focus of accurate object detection, but they contain the object that the CNN needs to locate. The layers of deep learning algorithms consist of stages to extract features in each layer and are very accurate.

The complexity of the image or shapes that this technique works on are examples of feature extraction by calculating the light or energy level. These layers are linked together in order to combine all feature extraction operations. Through the training phase, when positive features are present in the images the CNN gives these features more weight when they occur frequently, this will be considered. When a CNN is used to classify a specific image, it uses feature extraction, calculating weights based on the training phase, and finally compiling and analyzing the classification results. In addition, since some objects can be similar, images of various objects must also be provided for training. In order to calculate features, CNN performs calculations in iterations because the memory sizes of modern computers are insufficient to store all images [16-20]. Algorithms for cascaded feature detectors are based on feature filters, which are straightforward components that can be used to explain the alignment of particular pixels. Descriptive analysis is based on the component parts of the image, and revolves around their geometric positions in relation to each other. Whenever details and artifacts are detected at the appropriate geometric level, the model is adaptive because the algorithm may pretend that a specific feature exists even though the data contains more detail. In conclusion, neural network-based feature detectors may learn new patterns and be better able to adapt to difficult environments because they rely on self-learned features [21-24]. CNN's primary objective is to extract trouble-free features at a high resolution and transform them into more complex features at a lower resolution by subsampling a layer by a factor of two, which is a function of the convolution's kernel size. CNNs have been proposed for a long time, but their implementation complexities have made them unpopular in the engineering community. The constructed neural network receives input from the preferred facial appearances to

train the identifier on the seven common facial expressions. The CNN's objective is to train using a back-propagation algorithm that connects its input and output through a narrow conduit made of an invisible unit.

Because they are inextricably linked to the data sources, regularities cannot be separated from an unseen unit. Each system was designed to provide a maximum score of 1 for accurate facial recognition and a minimum score of 0 otherwise. Supposing the inputs in vectors to the network are denoted by; $X = [x_1, x_2, \dots, x_l]^T$, the output layers represented by; $Y = [y_1, y_2, \dots, y_k]^T$ with the prototypical for the optimization formulated as; $X: h \rightarrow Y$. The target datasets and its additive noise might respectively being represented as; $[t^1, t^2, \dots, t^k]$, with $u = [e^1, e^2, \dots, e^k]$, the eccentric K denotes the overall network patterns with the associated vectors represented as; $V = [v_1, v_2, \dots, v_{k-1}]$. The training algorithm is modelled with the following constraints when the training epochs are set to 1000 and the objective error is 0.001 [22-24].

$$\sum_{j=1}^k (t^j - y_2^j)^2 \quad (1)$$

For forward and backward propagation, fully connected CNN employ the following formulations:

$$\sum_i w_{j,i}^{T+1} x_i^T \quad (2)$$

$$\sum_i w_{j,i}^{T+1} g_i^T \quad (3)$$

Where $w_{j,i}^{T+1}$, g_i^T and x_i^T , are respectively, the weight connecting unit i at layer T to unit j at layer T+1 and the gradient and activation of unit I at layer T. The formulations above indicate that the activation units and gradients are pooled by the lower layer units connected to them and the higher layer units connected to them. However, due to the fact that connections leaving every piece are not constant lane of border effects, the pooling strategy becomes extremely painful and difficult to implement whenever calculating the gradients of the CNN. By dragging the incline from the upper layer, our model addresses this issue, and the CNN might be discriminately trained using the back propagation algorithm. A stochastic gradient decline might be utilized to update influences, as our formulation below demonstrates;

$$w_{i,j}(t) + \mu \frac{\partial C}{\partial w_{i,j}} \quad (4)$$

The cost function, which is heavily influenced by factors like the learning type and the activation function, is located between C and the learning rate. The activation function and cost function, respectively, are the SoftMax and cross entropy functions, and this study is being carried out using supervised learning on a multiclass classification problem. The following is a definition of the SoftMax function:

$$\frac{\exp(x_j)}{\sum \exp(x_k)} \quad (5)$$

The unit j is the output that indicates the class probability, and the total inputs to units k and j of the same level are presented by. The problem for multiclass classification in supervised learning is the cost function, which is well-defined as follows:

$$\sum_j d_j \log \log(P_j) - v_j \quad (6)$$

The weight is then linked to the preference unit in the unseen layer, and is the target probability for the output j. Weight is used to begin at the output unit, and error signals are transmitted value the input layer in a recursive fashion. As a cost function for the estimated error in the network's targeted function, the detected output compares favorably to the target valued, which was the facial image of the entire training set. The required minimization error is as follows:

$$\sum_{k=1}^K (t_k^T - g_k^T)^2 \quad (7)$$

such that, E^T is the resulting error, t_k^T , is the matching target value of the detected output (the activation of the k^{th} output unit) is where the error is well-defined, g_k^T [16-25].

2. Literature Review

In this part, sources and late examination connected with the subject of detecting counterfeit variety pictures in cloud networks will be surveyed. Significant important papers will be assessed to determine the commitment of analysts on this subject in the modern logical structure. A couple of compelling structures that have been made over the latest two or three years have embraced profound neural networks for face detection [2-8]. The face detection systems in view of profound neural networks might be partitioned towards single-stage too double stages moves close, as depicted as succeed: Lin Jiang 2021 [15], The eventual outcomes of relative examinations uncover that this innovation can unequivocally see different faces as well as outflanks speedy R-CNN to the extent that execution. In 2021, Prathmesh Deval, et. al., [16], fostering an area system of face shroud related with computerized clinical consideration administrations. By used OpenCV, to obtain induction to the live video move with furthermore for picture pre-handling. With face detection, Haar-Outpouring would become utilized, as it is an outstandingly strong face disclosure procedure. In 2021, S.Shivaprasad, et. al., [17] Profound learning, OpenCV, TensorFlow, additionally Keras have been used in the audit system to assist with the face acknowledgment utilizing covers. Through assisting with such a technique, security is ensured. The technique for face area has used the MobileNetV2 with CNN framework as a classifier; it is lightweight, which provides less boundaries, that could become used in installed contraptions (Onion Omega2 with Raspberry Pi) to perform genuine cover unmistakable proof. The precision of the approach utilized with such examination is 0.96; with the F1-score is 0.92. In 2021, Jansi Rani Sella Veluswami, et. al., [18], introduced a construction that was ready upon info of more than 11,000 pictures of faces either regardless of covers, utilizing different profound learning strategies. A SSDNET conspire is sent for face identification, and the outcome is perceived as a uniquely designed Lightweight CNN for cover area. Upon double unmistakable examining information, the design gets a striking precision of 96% with 96%. The plan would help authority organizations with prosperity leaders with battling the worldwide pandemic. In 2020, Guanhao Yang, et. al., [19], To apply YOLOV5, the best complaint revelation technique at the present time, in the genuine world, particularly toward wearing shroud in jam-pressed regions, it was recommended to substitute manual examination with simply a profound learning approach likewise use YOLOV5, the most impressive protest identification strategy at the second. Whenever visitors enter the shopping community, they will get pictures against the camera, which would straightaway, the shopping community entryway would be opened likewise shown to pass if face identified inside 2 sec is without a doubt a face have a cloak; more, it would become returned to face cover unmistakable proof even advancement is achieved. The test revelations suggest that the suggested computation could actually see face cover and additionally enable staff perception. In 2020, Hongtao, W., et. al., [20], redesigned the estimation to further foster proficiency of thing disclosure as the precision of the single-stage finder much of the time falls after the double stage indicator, the construction was ready upon VOC 2007 likewise 2012 models against all out for 16.551 pictures, of upgrading a segment of the info upset right as well however left as well as inconsistent testing may be used. The results show how the single-stage identifier obtains enormous accuracy on SSD. In 2019, Zhi Tian, et. al., [22], FCOS is a single-stage indicator that is without support and likewise proposition free. For examinations, FCOS outflanks ordinary support-according to single-stage identifiers like RetinaNet, Only Pull out all the stops, as well as SSD, yet with obviously less arrangement disarray. FCOS altogether avoids all anchor box assessment furthermore hyper-boundaries, addressing object revelation in a for every pixel estimate way, indistinguishable along those other thick Figure issues like semantic division. Between indicators, FCOS also shows up at cutting-edge adequacy. furthermore, show which FCOS could become used as RPNs in the double stage Quicker R-CNN indicator likewise surpasses its RPNs by means of an immense degree. In 2018, Qihang Wang, et. al., [21], suggested face distinguishing proof computation depending on Quick R-CNN that uses 3 estimations to find the candidate space of certain faces which could exist in the image (CNN framework, Haar-Ada-support procedure, as well as up-as well as-comer concentrate on heuristics). The contender region is dealt with into arranged convolution neural organization, which makes a last convolution quality (return for money invested) contingent upon Quick R-CNN network plan next a progression of

convolution likewise pooling systems. In 2018, Wenqi Wu, et. al., [16], introduced a face identifier using fluctuating records (DSFD) depending upon Quicker RCNN. The original organization could further develop facial acknowledgment accuracy while moreover going probably as a continuous Quicker R-CNN. All through FDDB info, 500 pictures were created heedlessly, the attempted distinction DSFD, standard methodology genuine Speedy R-CNN, MXnet, with UnitBox. The introduced system gets a 96.69% survey speed of 130.0 ms for handling picture outline, however genuine all the quicker R-CNN with MXnetutilise 140.0 ms likewise 230.0 ms, particularly, under the 700 misleading up-sides separate degree. In 2019, Q. Wang and J. Zheng, [24], introduced a face identifier using fluctuating metrics (DSFD) depending upon Quicker RCNN. The original organization could further develop facial acknowledgment accuracy while moreover going probably as a continuous Quicker R-CNN. The support is used against facial achievements to infer a human face idea. Then, by then, contingent upon the proposition metric, a Quick R-CNN for an equivalent sort network is introduced. By FDDB info, 500 pictures were created aimlessly, the attempted tests contrast DSFD, standard procedure genuine Quicker R-CNN, MXnet, and likewise UnitBox. The introduced procedure gets a 96.69% survey speed against 130.0 ms of handling a picture outline, however genuine Quicker R-CNN likewise MXnetutilise 140.0 ms with 230.0 ms, particularly, under the 700 misleading up-sides separate degree. In 2019, Zhu, K., et. al., [23], recommended Contextually Multi-Scaled Region-focused Convolutional Neural Networks (CMS-RCNNs) to determine different troublesome cases like critical facial hindrances, staggeringly lower objectives, extraordinary brightening, especially present changes, video or picture pressure goofs. Recommended networks, like region based CNNs, have a part of nearby proposition with a locale of interest (return for money invested) distinguishing parts. In any case, outside of those networks, two huge offers for further developing top tier facial acknowledgment execution.

3. Research Methods

To implement the proposed eye disease diagnosis and detection model, a set of data is used for eye medical diseases. The approved Internet sites provide many types of required data, and in this research, we relied on two important destinations for data processing, namely (kaggle.com and github.com), in addition to re-representing some of the data using the productive elements of the MATLAB program. A deep learning algorithm represented by fast convolutional neural networks (FCNNs) with initialization techniques and feature filtering and specifications has been proposed for use as a technique for analyzing, classifying and detecting eye diseases. The MATLAB environment was employed to design and test the proposed technique with the input data set.

3.1. The Implemented Dataset

The approved websites provide several types of required medical data. The two well-known websites were approved in this research to obtain the required data, namely (kaggle.com and github.com), where a large number of diseased eye image data were prepared with various classifications of known eye diseases. Figure 2 shows samples of the dataset that were used in the simulation.

By viewing the data models used, it can be seen that various types of eye diseases are combined in the applied data. The data set is uploaded by the prepared program, where the images are analyzed into a set of characteristics, according to the characteristics of each image, which represent a specific eye disease. The processes of classification and image analysis of the data set and the characterization of the characteristics of pathological cases will be clarified through the previous configuration and preparation processes of the proposed deep learning technique, as will be explained in the next article.

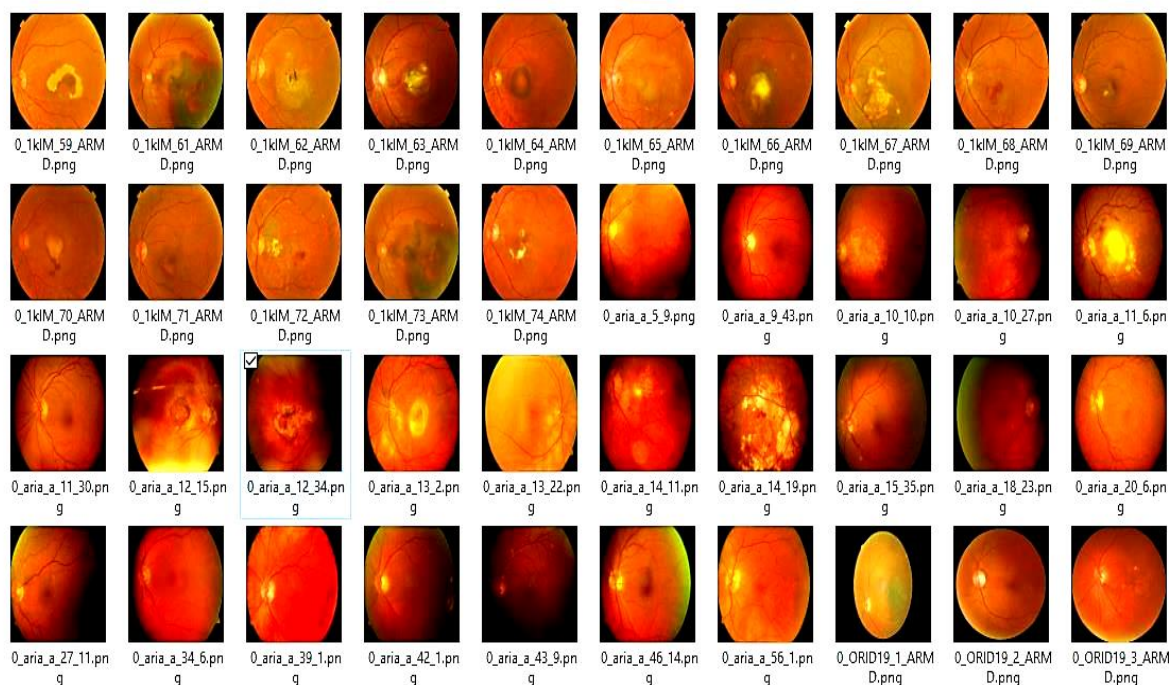
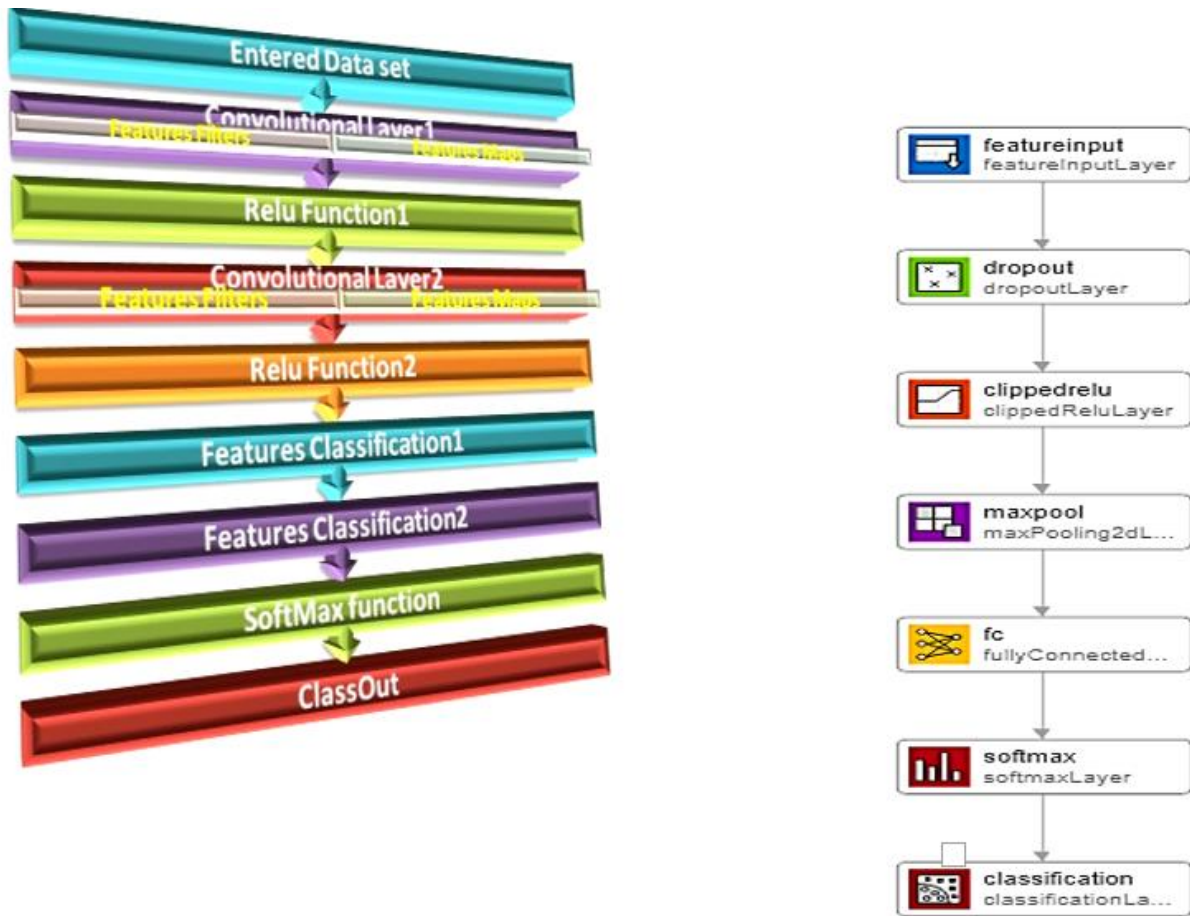


Figure 2.
Samples of the dataset which employed in the simulation.

3.2. The Suggested Eye Diseases Detection Model

In this Segment, the suggested model for detecting eye diseases will be introduced using deep learning FCNN algorithm. The flow chart of this first simulated F-CNN algorithm is introduced in Figure 3 (A).

By observing Figure 3 (B), the stream chart of the recommended F-CNN algorithm, clearly the stream chart of the F-CNN algorithm will made out of a few succeeded layers. The entered picture will initially go through the convolutional layer1 that will apply the elements filtering and mapping tasks. Then, the Relu function1 will decrease the weights of the data extricated from the convolutional layer1 to improve the component separated and planned data. Then, the resulting data will be gone through the convolutional layer2 for further element filtering and mapping. From that point forward, apply the Relu function2 to increase the data accuracy. The following layers in our proposed F-CNN algorithm are the elements classification1 and highlights classification2 which will address the regional classification of the CNN dissected data. The final stage in this proposed F-CNN algorithm is the SoftMax capability which will apply the final softening to the resulting data. The recently examined F-CNN proposed algorithm will be trained in our reenacted program using different training picture data sets containing every one of the clinical pictures of natural eye illnesses. The data set will be stacked from github.com and kaggle.com websites with many picture tests. Besides, the recommended F-CNN algorithm mimicked the model from entering data to the result results have displayed in Figure 3 (B), underneath.



(A) FR-CNN algorithm

(B) F-CNN algorithm

Figure 3.

(A) The flow chart of the suggested FR-CNN algorithm utilized for image classification model. (B) Structure of MATLAB simulated F-CNN algorithm.

In Figure 4, the introduced nitty gritty design of the proposed F-CNN algorithm reenacted model will develop from strong implicit capabilities operating to play out the profound learning technique to the entered datasets. Every one of the vital blocks, for example, featuring channels, convolutional layers, pooling capabilities, Relu capabilities, and classification utilities are accessible in this venture, the proposed model of the F-CNN algorithm model to play out the essential eye highlights detection and acknowledgment undertakings. Also, the detailed structure of the suggested FCNN layers is demonstrated in Figure 4.

```

>> TrainedNetwork.Layers

ans =

9x1 Layer array with layers:

   1 'imageinput'    Image Input          28x28x3 images with 'zerocenter' normalization
   2 'conv_1'        Convolution          32 3x3x3 convolutions with stride [1 1] and padding 'same'
   3 'relu_1'        ReLU                 ReLU
   4 'conv_2'        Convolution          20 3x3x32 convolutions with stride [1 1] and padding 'same'
   5 'relu_2'        ReLU                 ReLU
   6 'fc_1'          Fully Connected      395 fully connected layer
   7 'fc_2'          Fully Connected      2 fully connected layer
   8 'softmax'        Softmax              softmax
   9 'classoutput'   Classification Output crossentropyex with classes 'No_Tumor2' and 'Pituitary_Tumor2'

```

Figure 4.
The detailed structure of the suggested FCNN layers.

By following the Fig. above, we notice the details of the layers of the smart algorithm. It consists of an image input layer with a size of 28 x 28 x 3 images with “zero center” normalization. Followed by the 'conv_1' convolution layer of size 32 3x3x3 with step [1 1] and padding 'same'. It is then followed by the function 'relu_1' relu relu Followed by 'conv_2' convolution layer with size 20 3x3x32 convolution with step[1 1] and padding 'same' and then function 'relu_2' relu relu It is followed by the 'fc_1' fully connected layer 395 fully connected layers The 'fc_2' is fully connected layer 2 fully connected To be organized through the "softmax" function, softmaxsoftmax Finally, the “classoutput” classifies the crossentropyex output with the classes “No_Tumor2” and “Pituitary_Tumor2”. Moreover, the design parameters settings necessary to adjust and control the proposed FCNN algorithm have been outlined in Table 1.

Table 1.
The design specifications of the suggested FCNN algorithm model.

Initial learn rate	1e-6	1e-7	1e-9
Learn rate schedule	piecewise	piecewise	piecewise
Learn rate drop factor	0.1	0.15	0.2
Learn rate drop period	10	15	20
L2 regularization	1e-4	1e-5	1e-6
Max epochs	50	75	100

According to the design specifications illustrated in Table 1, the proposed FCNN algorithm will be employed using the MATLAB software which will be implemented to simulate the operation of the eye diseases features detection model.

4. Results and Discussions

In this section, we present the methodology for detecting eye diseases and identifying defect areas in input eye images using the proposed artificial intelligence deep learning algorithms, specifically Fully Convolutional Neural Networks (FCNN), implemented with MATLAB 2020b. The data was initialized, and its characteristics were filtered through layers complementary to the algorithm. Training involved over 500 images of diseased eyes to enable the smart network to recognize disease features and pinpoint defect areas within the input images. The program was executed according to the design specifications outlined in the methodology section, incorporating data images during the training phase of the algorithm. Figure 5 illustrates a sample test image of a human eye exhibiting a pathological defect, which is analyzed using the proposed FCNN technique.



Figure 5.
The test image sample of a human eye with a pathological defect examined through the proposed FCNN technique.

The training progress of the suggested FCNN algorithm is detailed in Figure 6, which showcases the efficiency levels relative to losses during both training and testing phases. The results indicate high training efficiency with minimal loss percentages. The pathological defect detections resulting from the FCNN algorithm are presented in Figure 7.

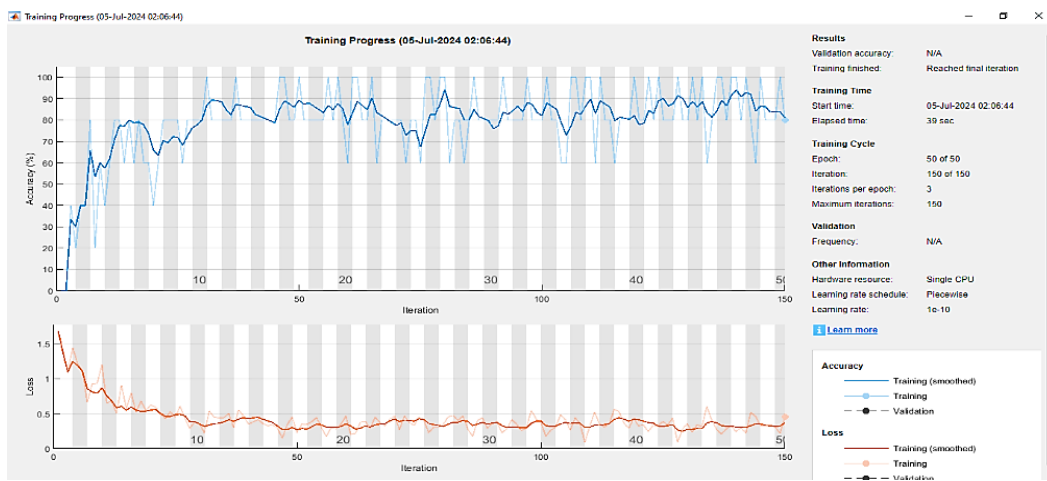


Figure 6.
The training progress of suggested deep learning FCNN algorithm.

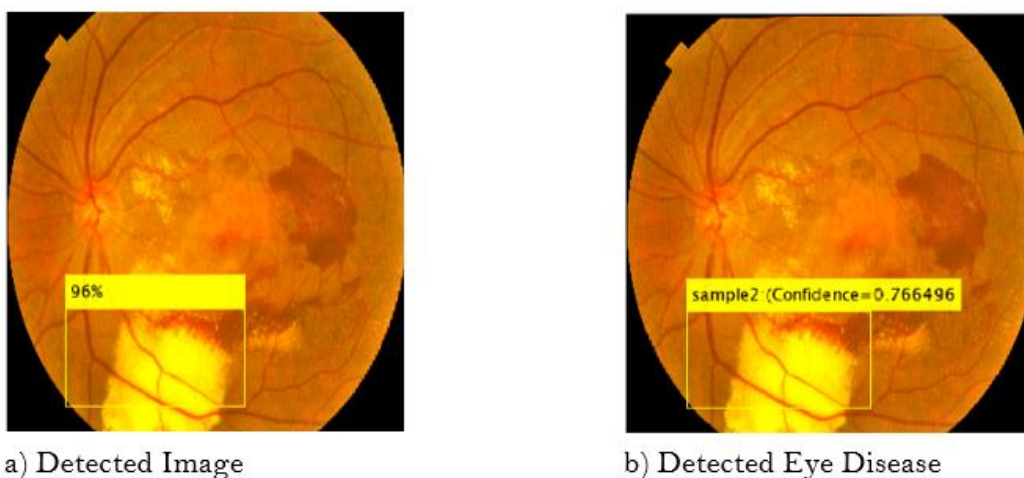


Figure 7.

The resulting pathological defect detection from the deep learning FCNN algorithm, (a) First detection, (b) Final recognition.

Figure 7 confirms the effectiveness of the proposed technique in identifying tumor and defect areas in the input human eye image, with rectangular boxes highlighting defect locations and focus areas. The training results reveal an impressive detection efficiency of 99-100% for identifying eye defects in the tested images. Further tests on additional image models of eye defects are illustrated in Figure 8.



Figure 8.

The second test image sample of a human eye with a pathological defect examined through the proposed FCNN technique.

Training progress for the second tested image is displayed in Figure 9, demonstrating similar trends in efficiency levels against losses, reaffirming the high performance of the training process.

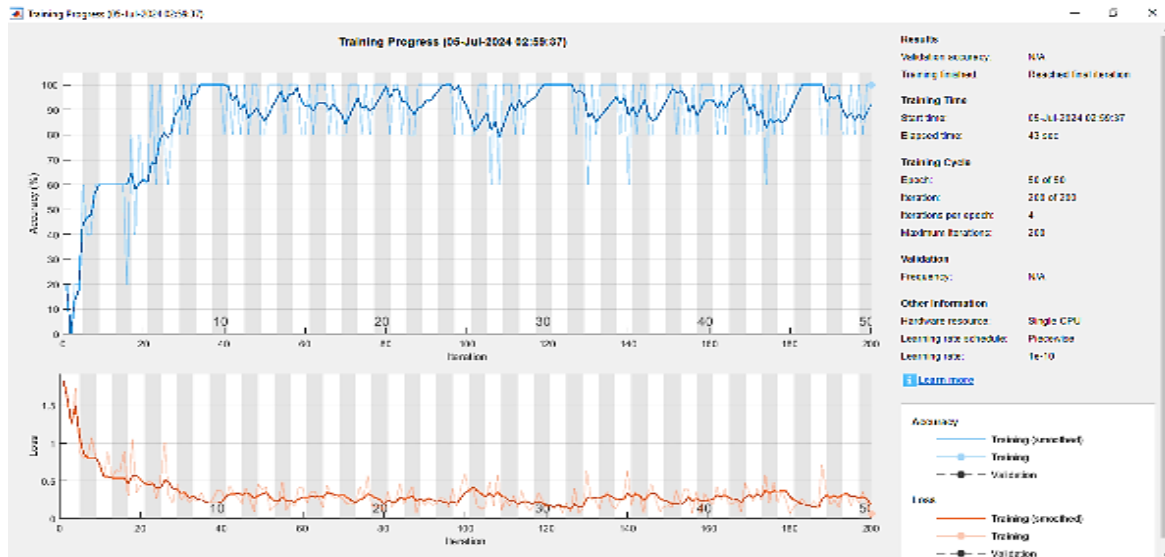


Figure 9.
The training progress of the suggested deep learning FCNN algorithm on the second tested image.

The resulting pathological defect detection for the second image is shown in Figure 10.

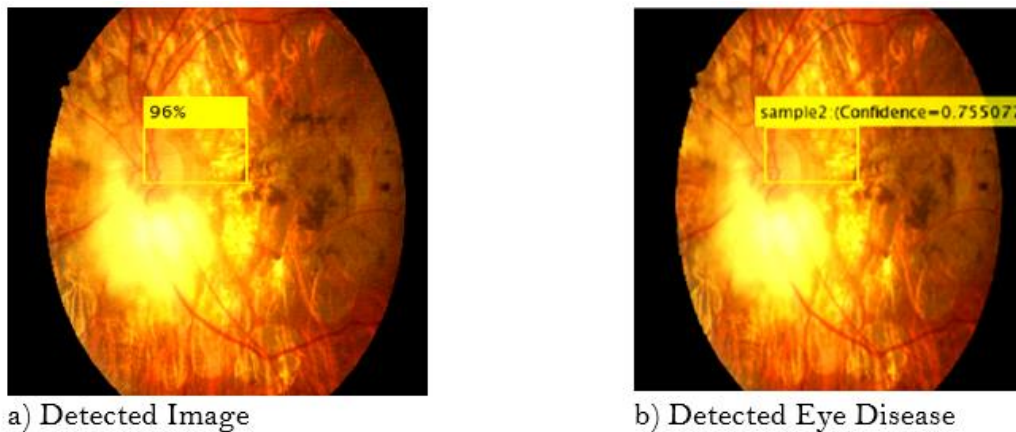


Figure 10.
The resulting pathological defect detection from the deep learning FCNN algorithm for the 2nd tested image, (a) First detection, (b) Final recognition.

The outcomes in Figure 10 further validate the proposed technique's capability to identify tumor and defect areas, marked by rectangular frames that indicate the location and level of focus. The findings consistently demonstrate efficiency values ranging from 99% to 100% in detecting eye defects across various input images.

This study highlights the successful application of fast convolutional neural networks for detecting human eye defects, based on the FCNN algorithm. The results indicate optimal detection performance across the tested eye image dataset. The deep learning FCNN approach was effectively applied to a wide range of input data models, encompassing initialization, analysis, training, and testing phases. Consequently, the detection efficiency was notably high, achieving accurate diagnoses of defect areas in eye images. Finally, an evaluation of the FCNN algorithm's performance metrics, as detailed in Table 2, reinforces the effectiveness of the proposed methodology.

Table 2.
The obtained FCNN evaluation metric records.

Metric values	TP	TN	FP	FN
	0.967	0.972	0.0125	0.0113
Accuracy	97.125%			
Specificity	98.5%			
Sensitivity	98.125%			
Precision	96.875			
F score	97.248%			

Based on the results in Table 2, it is concluded that the proposed FCNN algorithm for detecting eye defects has achieved a remarkable recognition efficiency of 99–100% across 50 examination epochs.

5. Conclusion

This research emphasizes the use of advanced deep learning algorithms, specifically Fully Convolutional Neural Networks (FCNN), for the classification of eye diseases within a diverse dataset of medical images. The training processes were meticulously designed, incorporating initialization mechanisms and the extraction of multi-specific features, which empowered the algorithm to effectively analyze input images, identify areas of concern, and accurately diagnose the locations of various conditions within the eye. Our results demonstrate a remarkable diagnostic efficiency of 99% and an error rate of no more than 0.015%, along with a training time of less than 35 seconds. These findings highlight the potential of deep learning techniques to enhance diagnostic accuracy and efficiency in ophthalmology.

Copyright:

© 2024 by the authors. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

References

- [1] Georgette A, Yaakov S, Christian H. Age-related disintegration in functional connectivity: Evidence from Reference Ability Neural Network (RANN) cohort. *Neuropsychologia*. 2021; 156:107856.
- [2] Fujimoto, S., van Hoof, H., Meger, D., 2018. Addressing function approximation error in actor-critic methods. arXiv preprint arXiv:1802.09477.
- [3] Haarnoja, T., Zhou, A., Abbeel, P., Levine, S., 2018. Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor. arXiv preprint arXiv:1801.01290.
- [4] Mohammed Jamal Mohammed, Jameel Kadhim Abed, Ali Jafer Mahdi, Basheer M. Hussein, "A Tuning of PID Power Controller using Particle Swarm Optimization for an Electro-Surgical Unit", *PRZEGLĄD ELEKTROTECHNICZNY*, ISSN 0033-2097, R. 99 NR 10/2023.
- [5] Huang QD, Dong XY, Chen DD, Zhou H, Zhang WM, Yu NH. Shape-invariant 3D adversarial point clouds. *IEEE Conference on Computer Vision and Pattern Recognition*. Vol. 18, 2022. p. 15314–23.
- [6] Ali Jafer Mahdi, Hussban Abood Saber, Ali Mohammed Ridha, Mohammed Jamal Mohammed, "COMPARATIVE ANALYSIS OF DIFFERENT INVERTERS AND CONTROLLERS TO INVESTIGATE PERFORMANCE OF ELECTROSURGICAL GENERATORS UNDER VARIABLE TISSUE IMPEDANCE", *Acta Innovations*, 2023, no. 48: 61–74, 61.
- [7] Irpan, A., 2018. Deep reinforcement learning doesn't work yet. Internet: <https://www.alexirpan.com/2018/02/14/rl-hard.html>.
- [8] Georgiou T, Liu Y, Chen W, Lew M. A survey of traditional and deep learning-based feature descriptors for high dimensional data in computer vision. *Int J Multimed Inf Retr*. 2019;9:135–70.
- [9] Lee, J., Won, J., Lee, J., Crowd simulation by deep reinforcement learning, *Proceedings of the 11th Annual International Conference on Motion, Interaction, and Games*. ACM, p. 2 (2018).
- [10] Li, X., Li, Q., Xu, X., Xu, D., Zhang, X., 2018. A novel approach to developing organized multi-speed evacuation plans. *Transactions in GIS* 22, 1205–1220.
- [11] Long, P., Fanl, T., Liao, X., Liu, W., Zhang, H., Pan, J., 2018. Towards optimally decentralized multi-robot collision avoidance via deep reinforcement learning, 2018 *IEEE International Conference on Robotics and Automation (ICRA)*. IEEE, pp. 6252–6259.

- [12] Lopes, G.C., Ferreira, M., da Silva Simões, A. and Colombini, E.L., 2018, November. Intelligent control of a quadrotor with proximal policy optimization reinforcement learning. In 2018 Latin American Robotic Symposium, 2018 Brazilian Symposium on Robotics (SBR) and 2018 Workshop on Robotics in Education (WRE), pp. 503-508.
- [13] Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., Killeen, T., Lin, Z., Gimselshein, N., Antiga, L. and Desmaison, A., 2019. PyTorch: An imperative style, high-performance deep learning library. In Advances in Neural Information Processing Systems (pp. 8024-8035).
- [14] Ali Mohammed Ridha, Ali Jafer Mahdi, Jameel Kadhim Abed, Shah Fahad, "PID fuzzy control applied to an electrosurgical unit for power regulation ", J Electr Bioimp, vol. 11, pp. 72-80, 2020 Received 28 Jul 2020 / published 15 Oct 2020 <https://doi.org/10.2478/joeb-2020-0011>
- [15] L. Jiang, J. Chen, H. Todo, Z. Tang, S. Liu and Y. Li, Application of a fast RCNN based on upper and lower layers in face recognition, Comput. Intell. Neurosci. 2021 (2021).
- [16] P. Deval, A. Chaudhari, R. Wagh, A. Auti and M. Parma, CNN based face mask detection integrated with digital hospital facilities, Int. J. Adv. Res. Sci., Commun. Tech. 4(2) (2021) 492-497.
- [17] S. Shivaprasad, M.D. Sai, U. Vignasahithi, G. Keerthi, S. Rishi and P. Jayanth, Real time CNN based detection of face mask using mobilenetv2 to prevent Covid-19, Ann. Roman. Soc. Cell Bio. 25(6) (2021) 12958-12969.
- [18] J.R. Sella Veluswami, S. Prakash and N. Parekh, Face mask detection using SSDNET and lightweight custom CNN, Proc. Int. Conf. IoT Based Control Networks and Intelligent Systems-ICICNIS 2021, (2021).
- [19] Ning Wang¹, Yang Gao¹, Hao Chen², Peng Wang¹, Zhi Tian², Chunhua Shen², Yanning Zhang¹, "NAS-FCOS: Fast Neural Architecture Search for Object Detection", arXiv:1906.04423v4 [cs.CV] 25 Feb 2020
- [20] W. Hongtao and Y. Xi, Object detection method based on improved one-stage detector, 5th Int. Conf. Smart Grid Electric. Autom. (ICSGEA), IEEE, 2020, pp. 209-212.
- [21] W. Wu, Y. Yin, X. Wang and D. Xu, Face detection with different scales based on faster R-CNN, IEEE Trans. Cyber. 49(11) (2018) 4017-4028.
- [22] Z. Tian, C. Shen, H. Chen and T. He, Fcos: Fully convolutional one-stage object detection, Proc. IEEE/CVF Int. Conf. Comput. Vision 2019, pp. 9627-9636.
- [23] Liu, J.; Zhu, K.; Lu, W.; Luo, X.; Zhao, X. A lightweight 3D convolutional neural network for deepfake detection. Int. J. Intell. Syst. 2021, 36, 4990-5004. [CrossRef]
- [24] Wang, Qihang, and Junbao Zheng. "Research on face detection based on fast R-CNN." Recent Developments in Intelligent Computing, Communication and Devices: Proceedings of ICCD 2017. Springer Singapore, 2019.
- [25] K. Lin, H. Zhao, J. Lv, C. Li, X. Liu, R. Chen and R. Zhao, Face detection and segmentation based on improved mask R-CNN, Discrete Dyn. Nature Soc. 2020 (2020).