

HMSANET: A HIERARCHICAL MULTI-SCALE ATTENTION NETWORK FOR PRECISE RETINAL LAYER SEGMENTATION IN COVID-19 OCT IMAGES

Radhwan A. A. Saleh^{1,2}, Lorena Álvarez-Rodríguez^{2,3}, Beatriz Cordon^{4,5}, Elena Garcia-Martin^{4,5}, Erchan Aptoula⁶, Joaquim de Moura^{2,3,}, Marcos Ortega^{2,3}*

¹CITIC, University of A Coruña, A Coruña, Spain

²VARPA Group, Biomedical Research Institute of A Coruña (INIBIC), University of A Coruña, A Coruña, Spain

³Department of Computer Science and Information Technologies, University of A Coruña, A Coruña, Spain

⁴Department of Ophthalmology, Miguel Servet University Hospital, Zaragoza, Spain

⁵Aragon Institute for Health Research (IIS Aragon), Miguel Servet Ophthalmology Innovation and Research Group (GIMSO), University of Zaragoza, Zaragoza, Spain

⁶Sabanci University, Faculty of Engineering and Natural Sciences (VPALab), Istanbul, Türkiye

ABSTRACT

The COVID-19 pandemic has underscored the systemic impact of the virus, including retinal abnormalities that may affect vision. Precise segmentation of retinal layers in optical coherence tomography (OCT) images is essential for reliable diagnosis and monitoring. We present the Hierarchical Multi-Scale Axial Attention Network (HMSANet), a novel deep learning architecture specifically designed for retinal layer segmentation in COVID-19 OCT data. HMSANet integrates two main modules: the Multi-Path Inception Attention Module (MPIAM) for multi-scale feature extraction, and the Multi-Scale Attention Module (MSAM) for adaptive attention weighting across spatial scales. This design enables the model to capture both fine-grained and large-scale retinal changes associated with COVID-19. Experimental results show that HMSANet consistently outperforms U-Net and SCAttNet by at least 1.2% in several retinal layers, demonstrating improved sensitivity to subtle structural alterations. By automating retinal layer segmentation, HMSANet enhances diagnostic speed and consistency, with potential applicability to other retinal diseases.

Index Terms— COVID-19, OCT, Retinal layer segmentation, Attention mechanisms, Deep learning

I. INTRODUCTION

Post-pandemic studies have shown that COVID-19 can affect organs beyond the respiratory system, with the retina receiving particular attention due to its structural vulnerability [1], [2]. Evidence suggests that infection may induce retinal alterations with potential impact on vision in recovered individuals. For example, recent findings report an increased likelihood of non-glaucomatous neuropathy and

greater choroidal thickness in patients with a history of COVID-19 [3].

A recent survey of 32 studies analyzing retinal anomalies in COVID-19 patients through OCT, OCT angiography, fundus photography, and autofluorescence imaging reported decreased vascular density, alterations in the retinal nerve fiber layer, and enlargement of the foveal avascular zone [4]. Fundus images often revealed cotton wool spots, microhemorrhages, and vascular occlusions, with severe cases showing higher prevalence of these findings. However, results varied due to patient demographics, comorbidities, and imaging timing. Despite methodological limitations and the absence of pre-infection retinal data, these studies emphasize the need for longitudinal research to clarify the long-term retinal effects of COVID-19.

Although manual OCT analysis has provided valuable insights, it is time-consuming, subjective, and insufficient for detecting subtle microstructural changes. The complexity of retinal OCT images, combined with resolution limitations of current devices, makes accurate identification of small alterations particularly challenging [5]. In addition, COVID-19-related retinal damage often introduces vascular anomalies and layer irregularities, further complicating segmentation. These challenges underscore the need for automated and highly accurate methods to assess retinal health in COVID-19 patients.

Deep learning-based segmentation has achieved notable success in retinal OCT analysis. However, conventional methods struggle with indistinct boundaries, noise, and disease-related anatomical alterations [6], [7]. Even state-of-the-art deep models face challenges in producing continuous and smooth retinal surfaces, particularly for thin layers affected by COVID-19-related microvascular changes [8].

Attention-based architectures have recently emerged as a promising solution, enhancing feature selection and improving segmentation accuracy in complex medical imaging tasks [9].

In this paper, we introduce HMSANet, a novel deep learning framework designed to analyze retinal changes associated with COVID-19. Unlike prior attention-based methods that often apply a single attention mechanism after multi-scale fusion, our approach introduces a hierarchical integration of two complementary attention modules. The key contributions of this work are: (1) The Multi-Path Inception Attention Module (MPIAM), which embeds independent channel attention within each parallel inception path, enabling scale-specific feature enhancement before aggregation; (2) The Multi-Scale Attention Module (MSAM), which facilitates attention-guided cross-scale interaction during feature fusion, allowing features at different resolutions to mutually inform each other; and (3) A unified architecture that jointly addresses scale-aware extraction and integration, demonstrating statistically significant improvements over state-of-the-art methods for OCT retinal segmentation.

II. METHODOLOGY

Fig. 1 provides an illustration of the suggested framework's approach. To establish a comprehensive baseline, we first curated a dedicated dataset of OCT images from COVID-19 patients and controls. This dataset, which is the first of its kind for studying the impact of COVID-19 on retinal structure, was meticulously annotated across ten anatomical layers by expert ophthalmologists. Following data preparation, the annotated OCT images are fed into HMSANet, a hierarchical encoder-decoder framework designed to capture both fine structural details and broader contextual information. The encoder progressively extracts multi-scale representations through a series of convolutional and pooling operations, while the decoder reconstructs high-resolution segmentation maps via skip connections and transposed convolutions. To enhance feature discrimination at different resolutions, MPIAM is embedded within encoder blocks to perform scale-aware feature extraction by selectively emphasizing informative channels across parallel receptive fields. Complementarily, MSAM is integrated into decoder blocks to facilitate attention-guided feature fusion, dynamically recalibrating concatenated multi-scale features and enabling effective cross-scale interaction during reconstruction. This hierarchical integration of multi-path and multi-scale attention allows HMSANet to adaptively balance local and global retinal cues, leading to improved sensitivity to subtle COVID-19-related retinal alterations.

Fig. 2 presents a visual illustration of the internal structures of the proposed attention modules. Both modules incorporate a channel-wise attention mechanism designed to adaptively recalibrate feature responses. Let $\mathbf{X} \in \mathbb{R}^{H \times W \times C}$ denote the input feature map, where H , W , and C represent

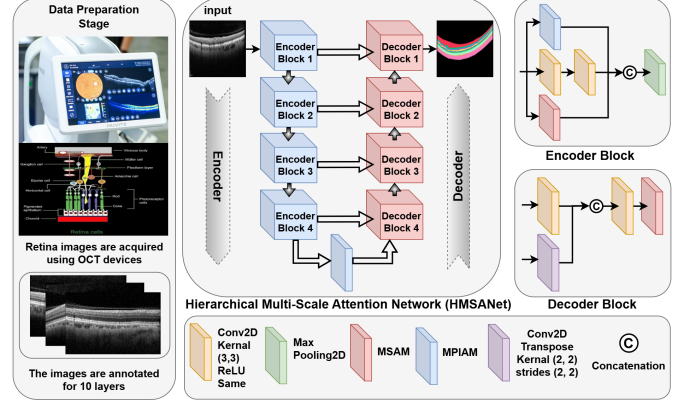


Fig. 1. Block diagram of the proposed HMSANet for automated retinal layer segmentation in COVID-19 patients.

the height, width, and number of channels, respectively. The attention mechanism first aggregates spatial information by applying global average pooling across the spatial dimensions:

$$\mathbf{z}_c = \frac{1}{HW} \sum_{i=1}^H \sum_{j=1}^W \mathbf{X}_{i,j,c}, \quad c = 1, 2, \dots, C, \quad (1)$$

yielding a channel descriptor $\mathbf{z} \in \mathbb{R}^C$.

This descriptor is then passed through a bottleneck transformation consisting of two fully connected layers with non-linear activation:

$$\mathbf{f} = \sigma(\mathbf{W}_2 \text{ReLU}(\mathbf{W}_1 \mathbf{z} + \mathbf{b}_1) + \mathbf{b}_2), \quad (2)$$

where $\mathbf{W}_1 \in \mathbb{R}^{\frac{C}{r} \times C}$ and $\mathbf{W}_2 \in \mathbb{R}^{C \times \frac{C}{r}}$ are learnable weights, r denotes the reduction ratio, $\mathbf{b}_1$ and $\mathbf{b}_2$ are bias terms, and $\sigma(\cdot)$ represents the sigmoid activation function.

The resulting attention vector is reshaped and applied to the input feature map via channel-wise multiplication:

$$\mathbf{Y} = \mathbf{X} \odot \text{Reshape}(\mathbf{f}), \quad (3)$$

where $\odot$ denotes element-wise multiplication.

II-A. Multi-Path Inception Attention Module (MPIAM)

The Multi-Path Inception Attention Module (MPIAM) enhances feature extraction by combining multi-scale receptive fields with scale-specific attention. Given an input feature map $\mathbf{X} \in \mathbb{R}^{H \times W \times C}$, MPIAM processes the input through four parallel paths, each capturing distinct spatial contexts.

In the first path, fine-grained local features are extracted using a 1×1 convolution followed by channel attention:

$$\mathbf{B}_1 = \mathcal{A}(\mathcal{C}_{1 \times 1}(\mathbf{X})), \quad (4)$$

where $\mathcal{C}_{k \times k}(\cdot)$ denotes a convolution with kernel size $k \times k$ and $\mathcal{A}(\cdot)$ represents the channel attention operation defined in Eq. (2).

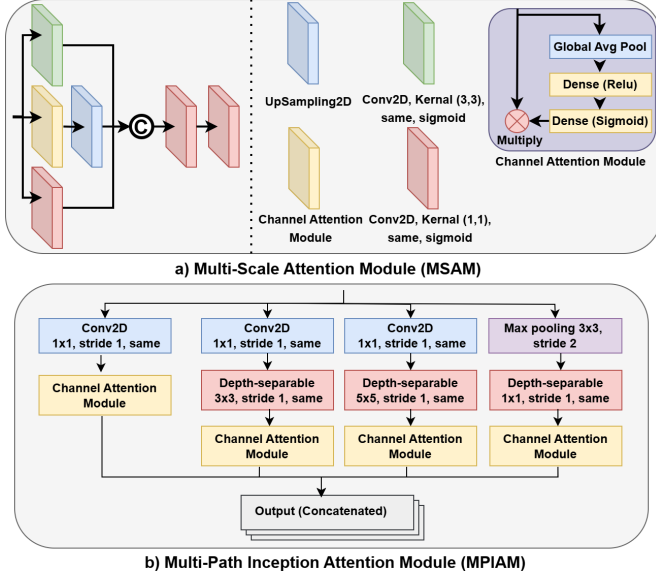


Fig. 2. Detailed structure of the proposed MSAM and MPIAM modules.

In the second path, mid-range contextual features are captured using a depthwise-separable 3×3 convolution:

$$\mathbf{B}_2 = \mathcal{A}(\mathcal{D}_{3 \times 3}(\mathcal{C}_{1 \times 1}(\mathbf{X}))), \quad (5)$$

In the third path, a larger receptive field is achieved using a depthwise-separable 5×5 convolution:

$$\mathbf{B}_3 = \mathcal{A}(\mathcal{D}_{5 \times 5}(\mathcal{C}_{1 \times 1}(\mathbf{X}))), \quad (6)$$

In the fourth path, global contextual information is emphasized using max pooling followed by a 1×1 convolution:

$$\mathbf{B}_4 = \mathcal{A}(\mathcal{C}_{1 \times 1}(\text{MaxPool}(\mathbf{X}))). \quad (7)$$

The outputs of all paths are concatenated along the channel dimension:

$$\mathbf{Y}_{\text{MPIAM}} = \text{Concat}(\mathbf{B}_1, \mathbf{B}_2, \mathbf{B}_3, \mathbf{B}_4), \quad (8)$$

producing a multi-scale representation with scale-specific channel recalibration.

II-B. Multi-Scale Attention Module (MSAM)

The Multi-Scale Attention Module (MSAM) is designed to refine feature representations by explicitly capturing contextual information at multiple spatial scales. Given an input feature map $\mathbf{X} \in \mathbb{R}^{H \times W \times C}$, where H , W , and C denote the height, width, and number of channels, respectively, MSAM extracts global, intermediate, and local features, which are subsequently fused and adaptively reweighted through an attention mechanism. To capture long-range dependencies and high-level semantic context, global channel-wise attention is first computed:

$$\mathbf{G} = \mathcal{A}(\mathbf{X}), \quad (9)$$

where $\mathcal{A}(\cdot)$ represents the channel attention operation defined in Eq. (2).

To capture spatial patterns at different resolutions, two parallel convolutional branches are employed. The intermediate-scale branch applies a standard 3×3 two-dimensional convolutional layer with stride 1 and same padding:

$$\mathbf{I} = \mathcal{C}_{3 \times 3}^{\text{C2D}}(\mathbf{X}; \mathbf{W}_I, \mathbf{b}_I), \quad (10)$$

which enlarges the receptive field and captures mid-level semantic information.

The local-scale branch applies a pointwise 1×1 two-dimensional convolutional layer with stride 1:

$$\mathbf{L} = \mathcal{C}_{1 \times 1}^{\text{C2D}}(\mathbf{X}; \mathbf{W}_L, \mathbf{b}_L), \quad (11)$$

which preserves spatial resolution while refining channel-wise interactions and fine-grained structural details.

The global, intermediate, and local feature maps are concatenated along the channel dimension:

$$\mathbf{F} = \text{Concat}(\mathbf{G}, \mathbf{I}, \mathbf{L}), \quad (12)$$

forming a unified multi-scale representation.

An adaptive attention map is then learned from the fused feature tensor using two successive pointwise 1×1 two-dimensional convolutional layers:

$$\mathbf{A} = \sigma(\mathcal{C}_{1 \times 1}^{\text{C2D}}(\sigma(\mathcal{C}_{1 \times 1}^{\text{C2D}}(\mathbf{F}; \mathbf{W}_1^a, \mathbf{b}_1^a)); \mathbf{W}_2^a, \mathbf{b}_2^a)), \quad (13)$$

where the sigmoid activation constrains attention values to the range $[0, 1]$.

Finally, the refined output of MSAM is obtained by modulating the input feature map using the learned attention map:

$$\mathbf{Y} = \mathbf{X} \odot \mathbf{A}, \quad (14)$$

where $\odot$ denotes element-wise multiplication. This operation enhances informative retinal structures while suppressing irrelevant responses, yielding a more discriminative representation for precise retinal layer segmentation.

III. EXPERIMENTAL SETUP

III-A. Dataset

The OCT images have been collected via the Heidelberg SPECTRALIS imaging platform, to compare the macular characteristics of patients with persistent and non-persistent COVID-19 symptoms against those of control (healthy) groups who had no prior history of the illness. The study included 10 participants for each group. Each individual in the control group has 25 uniformly spaced macular images while 61 images were retrieved from each patient in the other two groups. The images were preprocessed to standardize the region of interest in the macula, resulting in a uniform resolution of 512×434 pixels. Scans were included in

Table I. Mean performance metrics computed over 5-fold cross-validation with standard deviation in parentheses for each layer and model. Best results are highlighted in bold.

Layer	Model	IoU	F1 Score	Mean Surface Distance	Volumetric Similarity
ILM+RNFL	SCAttNet	0.8190 (± 0.0174)	0.9004 (± 0.0105)	6.0275 (± 0.1238)	0.9869 (± 0.0069)
	Unet	0.8807 (± 0.0115)	0.9365 (± 0.0065)	6.0819 (± 0.1370)	0.9904 (± 0.0046)
	HMSANet	0.8917 (± 0.0104)	0.9427 (± 0.0058)	6.0529 (± 0.1235)	0.9929 (± 0.0066)
GCL	SCAttNet	0.6906 (± 0.0236)	0.8168 (± 0.0165)	5.3013 (± 0.0893)	0.9823 (± 0.0094)
	Unet	0.7433 (± 0.0172)	0.8526 (± 0.0113)	5.3115 (± 0.1060)	0.9910 (± 0.0068)
	HMSANet	0.7575 (± 0.0154)	0.8619 (± 0.0099)	5.2252 (± 0.0403)	0.9936 (± 0.0035)
IPL	SCAttNet	0.7640 (± 0.0139)	0.8662 (± 0.0089)	5.3150 (± 0.0713)	0.9921 (± 0.0052)
	Unet	0.8114 (± 0.0106)	0.8958 (± 0.0064)	5.3115 (± 0.0301)	0.9886 (± 0.0055)
	HMSANet	0.8239 (± 0.0106)	0.9034 (± 0.0063)	5.3154 (± 0.0651)	0.9922 (± 0.0053)
INL	SCAttNet	0.7024 (± 0.0142)	0.8251 (± 0.0098)	5.1739 (± 0.0674)	0.9700 (± 0.0132)
	Unet	0.7699 (± 0.0125)	0.8699 (± 0.0080)	5.1074 (± 0.0512)	0.9880 (± 0.0089)
	HMSANet	0.7883 (± 0.0154)	0.8816 (± 0.0096)	5.0553 (± 0.1050)	0.9880 (± 0.0048)
OPL	SCAttNet	0.7343 (± 0.0121)	0.8467 (± 0.0081)	5.2205 (± 0.0457)	0.9703 (± 0.0071)
	Unet	0.8011 (± 0.0077)	0.8896 (± 0.0048)	5.1366 (± 0.0751)	0.9813 (± 0.0141)
	HMSANet	0.8106 (± 0.0047)	0.8954 (± 0.0028)	5.2309 (± 0.0718)	0.9849 (± 0.0056)
ONL	SCAttNet	0.8488 (± 0.0158)	0.9181 (± 0.0092)	5.4055 (± 0.0848)	0.9851 (± 0.0053)
	Unet	0.8908 (± 0.0087)	0.9422 (± 0.0049)	5.4387 (± 0.0600)	0.9898 (± 0.0032)
	HMSANet	0.8979 (± 0.0089)	0.9462 (± 0.0049)	5.4390 (± 0.0662)	0.9879 (± 0.0059)
ELM	SCAttNet	0.6802 (± 0.0214)	0.8095 (± 0.0151)	4.9898 (± 0.0423)	0.9826 (± 0.0207)
	Unet	0.7624 (± 0.0099)	0.8652 (± 0.0063)	4.9654 (± 0.0291)	0.9771 (± 0.0167)
	HMSANet	0.7744 (± 0.0095)	0.8729 (± 0.0060)	4.9388 (± 0.0538)	0.9868 (± 0.0132)
PL	SCAttNet	0.9019 (± 0.0121)	0.9484 (± 0.0067)	5.7702 (± 0.0543)	0.9866 (± 0.0119)
	Unet	0.9230 (± 0.0060)	0.9600 (± 0.0032)	5.7063 (± 0.0361)	0.9922 (± 0.0080)
	HMSANet	0.9277 (± 0.0074)	0.9625 (± 0.0040)	5.7197 (± 0.0384)	0.9896 (± 0.0071)
RPE	SCAttNet	0.6353 (± 0.0404)	0.7762 (± 0.0302)	5.0974 (± 0.1078)	0.9727 (± 0.0225)
	Unet	0.6626 (± 0.0297)	0.7966 (± 0.0217)	5.0648 (± 0.1051)	0.9811 (± 0.0066)
	HMSANet	0.6723 (± 0.0289)	0.8037 (± 0.0209)	5.1138 (± 0.0894)	0.9760 (± 0.0114)
Choroid	SCAttNet	0.7662 (± 0.0292)	0.8673 (± 0.0191)	7.4088 (± 0.0840)	0.9499 (± 0.0201)
	Unet	0.7871 (± 0.0151)	0.8808 (± 0.0094)	7.4305 (± 0.1279)	0.9610 (± 0.0183)
	HMSANet	0.8049 (± 0.0246)	0.8917 (± 0.0152)	7.4813 (± 0.1512)	0.9756 (± 0.0156)

the dataset only if they had a quality rating of at least 25/40 points, as determined by a thorough quality evaluation [10]. The dataset was then annotated using CVAT. The participants in the dataset were well-matched in terms of age and gender distribution to reduce potential biases that could affect the study’s results. The inclusion criteria and data collection techniques were carefully planned to ensure representativeness and reliability.¹

III-B. Training details

The Nvidia GeForce RTX 4060 GPU was used in our experiments. Data was divided using patient-level partitioning and evaluated with 5-fold cross-validation. Each fold was used successively as a validation set, with the remaining four folds serving as the training set. The images and their corresponding masks were normalized and resized to

256×256 . The batch size used was 1, the learning rate is 1×10^{-4} , the Adam algorithm was used for optimization, and the number of epochs is 50 for each run. The loss function is a combination of Dice loss and Focal loss.

III-C. Validation and state of the art comparison

Validation metrics used are Intersection over Union (IoU), F1 Score, Mean Surface Distance, and Volumetric Similarity. To provide a comprehensive comparison, we selected U-Net [11] and SCAttNet [12] as baseline models. These represent a well-established standard architecture and a state-of-the-art attention-based method for medical segmentation, respectively, providing a strong baseline for evaluating the proposed improvements.

IV. RESULTS AND DISCUSSION

The dataset was partitioned at the patient level and evaluated using five-fold cross-validation to ensure robust and unbiased performance assessment. This strategy prevents data

¹The Ethics Committee of Miguel Servet Hospital approved this investigation with registration number C.I. PI21/113.

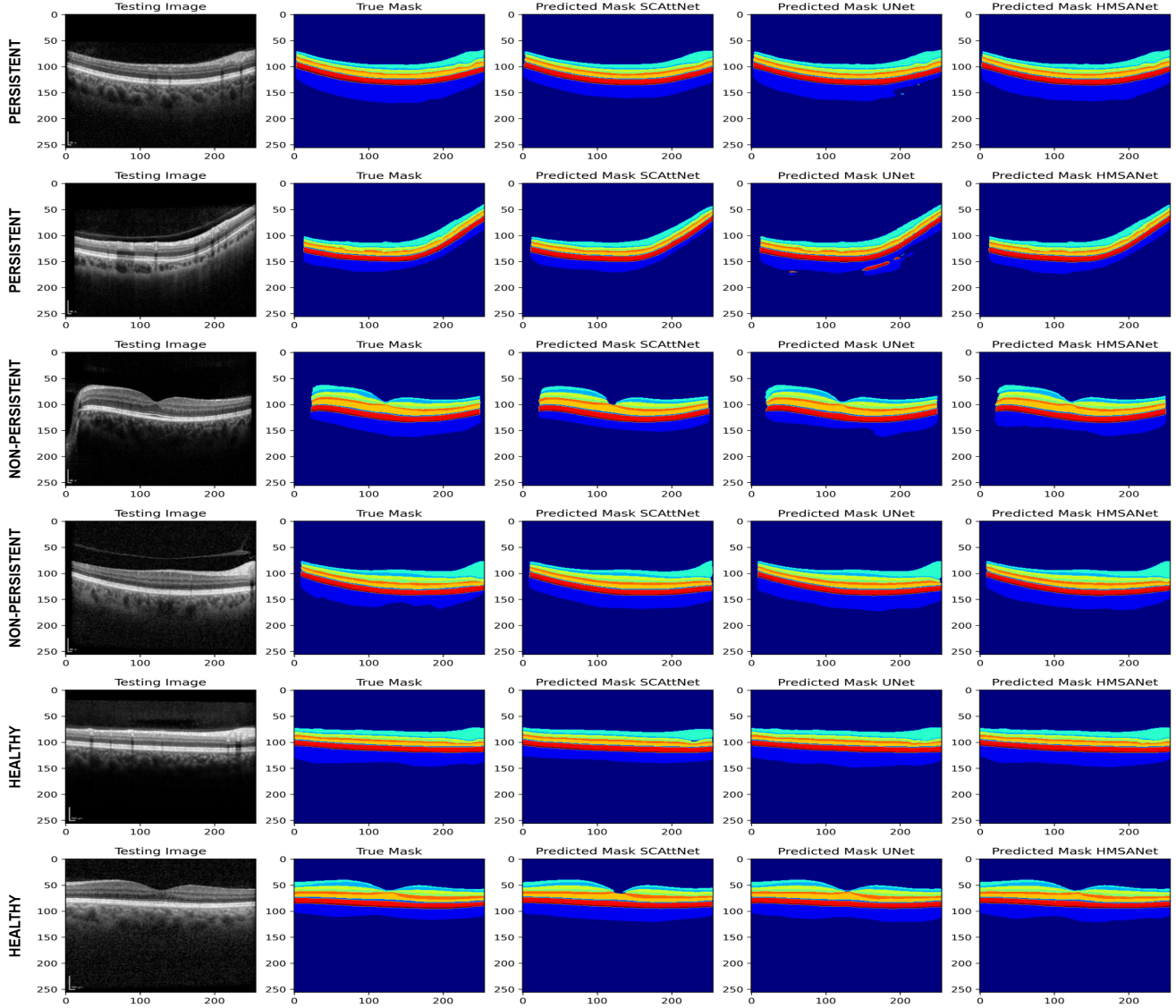


Fig. 3. Retinal layer segmentation in OCT images using SCAttNet, UNet, and proposed HMSANet across different conditions.

leakage between training and validation sets and provides a reliable estimate of model generalization across subjects.

Table II. Statistical comparison of model performance across 10 retinal layers using Wilcoxon signed-rank tests.

Comparison	Median Diff	W-statistic	p-value	Effect Size (r)
HMSANet vs U-Net	0.0115	55	9.77e-04	0.886
HMSANet vs SCAttNet	0.0634	55	9.77e-04	0.886

Quantitative results are summarized in Table I. The proposed HMSANet consistently outperforms both SCAttNet [12] and U-Net [11] across all evaluation metrics,

including Intersection over Union (IoU), F1 Score, Mean Surface Distance, and Volumetric Similarity. This performance gain is observed across all retinal layers, demonstrating the robustness of the proposed architecture for comprehensive retinal layer segmentation. Notably, HMSANet exhibits superior performance in the outer retinal layers, which are clinically significant for identifying pathological changes associated with COVID-19. Improvements in these layers indicate the model's enhanced ability to capture subtle structural variations, which can be attributed to the effective integration of multi-path feature extraction and adaptive attention mechanisms. Consistent gains are also observed in the inner retinal layers, such as the ILM+RNFL and GCL, highlighting HMSANet's ability to handle fine-grained

Table III. Per-layer segmentation performance metrics for HMSANet and its variants. Best results are highlighted in bold.

Layer	Model	IoU	F1 Score	Mean Surface Distance	Volumetric Similarity
ILM+RNFL	HMSANet w/o MPIAM	0.8857 (± 0.0121)	0.9393 (± 0.0068)	6.100 (± 0.2251)	0.9867 (± 0.0096)
	HMSANet w/o MSAM	0.8862 (± 0.0141)	0.9396 (± 0.0079)	6.170 (± 0.2299)	0.9921 (± 0.0036)
	HMSANet	0.8917 (± 0.0116)	0.9427 (± 0.0065)	6.053 (± 0.1400)	0.9929 (± 0.0073)
GCL	HMSANet w/o MPIAM	0.7488 (± 0.0187)	0.8563 (± 0.0122)	5.224 (± 0.0800)	0.9828 (± 0.0092)
	HMSANet w/o MSAM	0.7466 (± 0.0174)	0.8548 (± 0.0113)	5.240 (± 0.1521)	0.9825 (± 0.0137)
	HMSANet	0.7575 (± 0.0172)	0.8619 (± 0.0111)	5.230 (± 0.0451)	0.9936 (± 0.0039)
IPL	HMSANet w/o MPIAM	0.7868 (± 0.0154)	0.8806 (± 0.0097)	5.110 (± 0.0986)	0.9868 (± 0.0047)
	HMSANet w/o MSAM	0.7833 (± 0.0160)	0.8784 (± 0.0100)	5.190 (± 0.0868)	0.9891 (± 0.0078)
	HMSANet	0.7883 (± 0.0172)	0.8816 (± 0.0107)	5.055 (± 0.1200)	0.9880 (± 0.0053)
INL	HMSANet w/o MPIAM	0.8184 (± 0.0115)	0.9001 (± 0.0070)	5.400 (± 0.1007)	0.9913 (± 0.0065)
	HMSANet w/o MSAM	0.8188 (± 0.0110)	0.9004 (± 0.0067)	5.330 (± 0.0496)	0.9910 (± 0.0086)
	HMSANet	0.8239 (± 0.0118)	0.9034 (± 0.0071)	5.315 (± 0.0700)	0.9922 (± 0.0059)
OPL	HMSANet w/o MPIAM	0.8953 (± 0.0108)	0.9447 (± 0.0060)	5.433 (± 0.0800)	0.9941 (± 0.0038)
	HMSANet w/o MSAM	0.8963 (± 0.0103)	0.9453 (± 0.0057)	5.500 (± 0.0773)	0.9936 (± 0.0036)
	HMSANet	0.8979 (± 0.0099)	0.9462 (± 0.0055)	5.440 (± 0.0740)	0.9879 (± 0.0066)
ONL	HMSANet w/o MPIAM	0.8059 (± 0.0080)	0.8925 (± 0.0049)	5.210 (± 0.0715)	0.9859 (± 0.0128)
	HMSANet w/o MSAM	0.8060 (± 0.0082)	0.8926 (± 0.0050)	5.192 (± 0.0700)	0.9862 (± 0.0109)
	HMSANet	0.8106 (± 0.0052)	0.8954 (± 0.0032)	5.230 (± 0.0802)	0.9849 (± 0.0063)
ELM	HMSANet w/o MPIAM	0.7684 (± 0.0143)	0.8689 (± 0.0091)	4.906 (± 0.0700)	0.9815 (± 0.0203)
	HMSANet w/o MSAM	0.7699 (± 0.0143)	0.8699 (± 0.0090)	4.970 (± 0.0902)	0.9893 (± 0.0130)
	HMSANet	0.7744 (± 0.0106)	0.8729 (± 0.0067)	4.940 (± 0.0602)	0.9868 (± 0.0148)
PL	HMSANet w/o MPIAM	0.9269 (± 0.0093)	0.9620 (± 0.0050)	5.692 (± 0.1100)	0.9915 (± 0.0068)
	HMSANet w/o MSAM	0.9259 (± 0.0069)	0.9615 (± 0.0037)	5.760 (± 0.0855)	0.9897 (± 0.0070)
	HMSANet	0.9277 (± 0.0083)	0.9625 (± 0.0044)	5.720 (± 0.0430)	0.9896 (± 0.0079)
RPE	HMSANet w/o MPIAM	0.6694 (± 0.0367)	0.8015 (± 0.0265)	5.042 (± 0.1100)	0.9746 (± 0.0131)
	HMSANet w/o MSAM	0.5279 (± 0.2957)	0.6359 (± 0.3558)	5.130 (± 0.0760)	0.7805 (± 0.4364)
	HMSANet	0.6723 (± 0.0324)	0.8037 (± 0.0234)	5.110 (± 0.0999)	0.9760 (± 0.0127)
Choroid	HMSANet w/o MPIAM	0.8029 (± 0.0242)	0.8905 (± 0.0149)	7.360 (± 0.2453)	0.9664 (± 0.0131)
	HMSANet w/o MSAM	0.8043 (± 0.0189)	0.8915 (± 0.0116)	7.333 (± 0.1700)	0.9699 (± 0.0127)
	HMSANet	0.8049 (± 0.0275)	0.8917 (± 0.0170)	7.480 (± 0.1690)	0.9756 (± 0.0175)

anatomical structures that are typically challenging due to low contrast and thin boundaries.

Segmenting inner retinal layers remains particularly difficult because of their complex morphology and close spatial proximity. As shown in Table I, HMSANet achieves improved segmentation accuracy across these layers compared to baseline models. These improvements can be largely attributed to the Multi-Path Inception Attention Module (MPIAM), which enables parallel extraction of features at different receptive fields, thereby enhancing boundary delineation and preserving structural continuity.

The Retinal Pigment Epithelium (RPE) layer presents a further challenge due to its thin structure and variable appearance. Despite this complexity, HMSANet achieves the best overall performance among the evaluated methods, indicating its ability to maintain segmentation fidelity in both superficial and deeper retinal layers. This suggests that the proposed attention mechanisms effectively balance global contextual understanding with local structural refinement.

To assess the statistical significance of these improvements, we performed Wilcoxon signed-rank tests comparing HMSANet against the other two models using the IoU values across the 10 retinal layers. This non-parametric test was chosen due to the small sample size ($n = 10$) and the potential lack of independence between adjacent retinal layers, which may exhibit spatial correlation due to their anatomical ordering. Table II presents the statistical comparison results. The analysis reveals that HMSANet achieves statistically significant improvements over both U-Net ($p = 9.77\text{e-}04$) and SCAttNet ($p = 9.77\text{e-}04$). The large effect sizes ($r = 0.886$ for both comparisons) indicate substantial practical significance, confirming that the observed performance gains are both statistically and clinically meaningful.

Qualitative segmentation results are illustrated in Fig. 3, which compares predictions from SCAttNet, U-Net, and HMSANet across persistent, non-persistent, and healthy cases. HMSANet produces segmentation masks that most closely align with the ground truth, particularly in patho-

logical cases where retinal layer boundaries are disrupted. In contrast, U-Net exhibits minor artifacts and boundary inaccuracies, while SCAttNet, although competitive, shows inconsistencies in fine structural details. These visual results corroborate the quantitative findings and demonstrate the practical advantages of HMSANet in real-world clinical scenarios.

Table IV. Ablation study statistical comparison across 10 retinal layers using Wilcoxon signed-rank tests.

Comparison	Median Diff	W-statistic	p-value	Effect Size (r)
HMSANet vs w/o MPIAM	0.0045	39	0.1377	0.371
HMSANet vs w/o MSAM	0.0050	41	0.0967	0.435

An ablation study was conducted to assess the individual contributions of the proposed MPIAM and MSAM modules. The results are reported in Table III. To evaluate the significance of each module’s contribution, we performed Wilcoxon signed-rank tests comparing the full HMSANet against its variants (w/o MPIAM and w/o MSAM) using the IoU values across the 10 retinal layers. Similar to the main comparison, this non-parametric approach was selected to avoid assumptions of normality and independence, given the small sample size and anatomical ordering of layers. Table IV presents the statistical comparison results. The analysis reveals that removing the MPIAM module leads to a performance degradation (median difference = 0.0045), though the result did not reach statistical significance ($p = 0.1377$, effect size $r = 0.371$). Similarly, the removal of the MSAM module yielded a median difference of 0.0050 with $p = 0.0967$ and effect size $r = 0.435$, also falling short of statistical significance. While these results do not confirm a statistically significant contribution of each individual module in isolation, the consistent performance improvements observed across individual layers in Table III, combined with the complementary behavior discussed below, suggest that both modules contribute meaningfully to the overall architecture, particularly in a synergistic manner.

The ablation study further revealed interesting complementary behavior between the two modules. For certain layers (e.g., ONL and PL), the removal of MSAM had a more pronounced effect than the removal of MPIAM, suggesting that MSAM plays a critical role in refining boundaries for specific anatomical structures. Conversely, for layers with complex textural patterns (e.g., RPE and Choroid), MPIAM contributed more substantially, likely due to its ability to capture multi-scale contextual information. This complementary behavior underscores the synergy achieved by integrating both attention mechanisms within a unified architecture.

V. CONCLUSION

This study addressed the critical need for precise, automated segmentation of retinal layers in OCT scans from COVID-19 patients, a challenging problem due to subtle, disease-induced anatomical changes that compromise traditional and deep learning methods. To solve this, we introduced HMSANet, a framework that integrates multi-path and multi-scale attention mechanisms (MPIAM and MSAM) to enhance feature extraction and fusion across varying receptive fields. The core innovation lies in its hierarchical integration of these two dedicated attention mechanisms, where MPIAM tackles the problem of multi-scale feature extraction by employing parallel paths with scale-specific channel attention. This allows the model to discern fine-grained details crucial for thin layers alongside broader contextual information. While MSAM addresses the challenge of feature fusion by dynamically recalibrating and integrating information across different resolutions, ensuring that the segmentation remains coherent and continuous even when layer textures are disrupted.

Our evaluation on the collected OCT dataset from COVID-19 patients and controls demonstrates that HMSANet successfully solves the outlined problem. It consistently surpasses leading benchmarks like U-Net and SCAttNet across all retinal layers and all quantitative metrics (IoU, F1 Score, Mean Surface Distance, Volumetric Similarity). Qualitatively, it produces smoother, more accurate layer boundaries that closely follow the ground truth, especially in pathological cases where other models falter. The ablation study confirms that both proposed modules are essential to this performance gain. In summary, HMSANet provides an effective and reliable automated solution for quantifying post-COVID retinal changes. By enhancing diagnostic consistency and reducing reliance on manual analysis, it offers significant clinical utility for the early detection and longitudinal monitoring of COVID-19 ocular sequelae, serving as a valuable tool for both clinicians and researchers.

ACKNOWLEDGEMENT

This study was supported by the Carlos III Health Institute through grants PI23/00935 and RD21/0002/0050 (Inflammatory Disease Network – RICORS), and by the Government of Aragon through group B23_23R. It also received funding from the Ministerio de Ciencia e Innovación (PID2023-148913OB-I00) and from the Consellería de Educación, Universidade e Formación Profesional, Xunta de Galicia, under the Grupos de Referencia Competitiva program (ED431C 2024/33). Additional support was provided by the Carlos III Health Institute through grant FORT23/00010 (FORTALECE Program), the European Union via the 3i ICT project (H2020-MSCA-COFUND-2020, grant agreement GA 101034261), and the Consellería de Cultura, Educación, Formación Profesional e Universidades of the Xunta de Galicia through the 3i ICT COFUND agreement.

”Consellería de Educación, Ciencia, Universidades e Formación Profesional (Xunta de Galicia - Convenio para o desenvolvemento de accións estratéxicas de I+D+i 2025-2026.

ETHICS APPROVAL

This observational study was conducted in accordance with the Declaration of Helsinki and approved by the Ethics Committee of Investigation of Miguel Servet Hospital (CE-ICA) (C.I. PI21/113).

VI. REFERENCES

- [1] Hans E. Grossniklaus, Eldon E. Geisert, and John M. Nickerson, “Introduction to the retina,” *Progress in Molecular Biology and Translational Science*, vol. 134, pp. 383–396, 2015.
- [2] Mehmet Selim Gel, Ayhan Kanat, Demet Seker, Hakan Koc, Iskender Samet Daltaban, Huseyin Findik, and Omer Lutfi Gundogdu, “Changes in retinal nerve fiber layer thickness may be the cause of post-covid-19 headaches,” *Neurological Research*, pp. 1–10, 2024.
- [3] Ilda Maria Poças, Pedro Lino, Carina Silva, Paula Mendonça, João Paulo Cunha, Olga Barroqueiro, Francisca Carvalho, Inês Nicho, Mariana Castelhana, Patrícia Condado, et al., “Ocular repercussions in covid-19 patients: structural changes of the retina and choroid,” *Strabismus*, vol. 31, no. 4, pp. 271–280, 2023.
- [4] Orlaith E. McGrath and Tariq M. Aslam, “Use of imaging technology to assess the effect of covid-19 on retinal tissues: A systematic review,” *Ophthalmology and Therapy*, vol. 11, no. 3, pp. 1017–1030, 2022.
- [5] Yuhe Shen, Jiang Li, Weifang Zhu, Kai Yu, Meng Wang, Yuanyuan Peng, Yi Zhou, Liling Guan, and Xinjian Chen, “Graph attention u-net for retinal layer surface detection and choroid neovascularization segmentation in oct images,” *IEEE Transactions on Medical Imaging*, vol. 42, no. 11, pp. 3140–3154, 2023.
- [6] Reza Rasti, Armin Biglari, Mohammad Rezapourian, Ziyun Yang, and Sina Farsiu, “Retifluidnet: A self-adaptive and multi-attention deep convolutional network for retinal oct fluid segmentation,” *IEEE Transactions on Medical Imaging*, vol. 42, no. 5, pp. 1413–1423, 2022.
- [7] Guogang Cao, Yan Wu, Zeyu Peng, Zhilin Zhou, and Cuixia Dai, “Self-attention cnn for retinal layer segmentation in oct,” *Biomedical Optics Express*, vol. 15, no. 3, pp. 1605–1617, 2024.
- [8] Qifeng Yan, Yuhui Ma, Wenjun Wu, Lei Mou, Wei Huang, Jun Cheng, and Yitian Zhao, “Choroidal layer analysis in oct images via ambiguous boundary-aware attention,” *Computers in Biology and Medicine*, vol. 175, pp. 108386, 2024.
- [9] Bingbing Li, Yao Cong, and Hongwei Mo, “Ophthalmic oct segmentation method based on rnn-attention,” *IEEE Access*, vol. 11, pp. 129601–129612, 2023.
- [10] Axel Petzold, Philipp Albrecht, Laura Balcer, Erik Bekkers, Alexander U. Brandt, Peter A. Calabresi, Orla Galvin Deborah, Jennifer S. Graves, Ari Green, Pearse A. Keane, et al., “Artificial intelligence extension of the oscar-ib criteria,” *Annals of Clinical and Translational Neurology*, vol. 8, no. 7, pp. 1528–1542, 2021.
- [11] Olaf Ronneberger, Philipp Fischer, and Thomas Brox, “U-net: Convolutional networks for biomedical image segmentation,” in *Medical Image Computing and Computer-Assisted Intervention*, Nassir Navab, Joachim Hornegger, William M. Wells, and Alejandro F. Frangi, Eds. 2015, vol. 9351 of *Lecture Notes in Computer Science*, pp. 234–241, Springer.
- [12] Haifeng Li, Kaijian Qiu, Li Chen, Xiaoming Mei, Liang Hong, and Chao Tao, “Scattnet: Semantic segmentation network with spatial and channel attention mechanism for high-resolution remote sensing images,” *IEEE Geoscience and Remote Sensing Letters*, vol. 18, no. 5, pp. 905–909, 2020.