

CONTINUAL LEARNING-DRIVEN LAPLACIAN PALE TRANSFORMATIVE CONVOLUTION NETWORK FRAMEWORK FOR DIABETIC RETINOPATHY DETECTION FROM RETINAL FUNDUS IMAGES

K.S.NALINI^{1*}, Dr.ARUNACHALAM A.S.²

^{1*}Research scholar, Department of Computer Science, Vel's Institute of Science and Technology and Advanced Studies (VISTAS), Tamil Nadu, India.

²Professor, Department of Computer Science, Vel's Institute of Science and Technology and Advanced Studies (VISTAS), Tamil Nadu, India.

Corresponding Author: ksnalini89@gmail.com

ABSTRACT

Diabetic retinopathy (DR) is a diabetes-related eye ailment caused by retinal blood vessel (BV) damage. This manuscript presents a novel continual DL framework for efficient DR disease detection. Initially, input retinal fundus images are taken from the IDRI Dataset for accurate DR disease detection, which undergoes the preprocessing stage by employing a Color Wiener Filter (CWF) that can enhance image clarity by adaptively removing noise while maintaining edge details for further processing. After preprocessing, a novel Laplacian Pale Transformative Convolution network (LPTCN) is introduced, which classifies with more accuracy the distinction between different DR abnormalities. Moreover, the proposed framework integrates Elastic Weight Consolidation (EWC) and Herding Selection Replay (HSR) to prevent catastrophic forgetting on new data samples. The proposed framework is simulated in the Python platform. In the simulation part, the average Accuracy of 97%, Matthew's Correlation Coefficient (MCC) of 0.936, symmetric mean absolute percentage error (SMAPE) of 2.07, and Computation Time (CT) of 4.9s, Youden's index (YI) of 0.89 are obtained by the suggested framework on DR identification.

Keywords: *Diabetic Retinopathy (DR), Retinal Fundus Image (RFI), Multi-class Classification, Image Preprocessing, Continual Learning, Laplacian Pale Transformative Convolution Network.*

1. INTRODUCTION

Diabetic Retinopathy (DR) is a severe microvascular diabetes mellitus complication that may result in irreversible vision impairment and blindness, especially in working-age individuals. DR destroys the retinal blood vessels (BV), which provide oxygen and nutrients to the retina [1]. DR is classified into two categories: Non-Proliferative DR (NPDR) and Proliferative DR (PDR). NPDR is also divided into mild, moderate, and severe stages, and PDR is the final stage. DR can also be classified on the basis of severity into five phases, from Phase 0 (no DR) to Phase 4 (PDR). DR is responsible for causing almost 2.6% of blindness in the whole world, with more than 239 million sufferers in 2010 [2, 3]. The International Diabetes Federation (IDF) noted that as of 2017, 449 million people had diabetes, and over a third had signs of DR. This is estimated to increase to

592 million by 2025, reflecting the expanding worldwide burden.

Research Problem: DR needs early detection because the disease is usually symptom-free in the initial stages, and floaters, blurred vision, and loss of vision become manifest only at late stages [4]. Although color fundus imaging is common for DR diagnosis, manual grading is time-consuming, subjective, and susceptible to inter-observer variability. Machine learning (ML) and deep learning (DL) based automated classification systems have been promising, but challenges, including high computational expense, low availability of annotated datasets, overfitting, and bias, limit their implementation in real-world clinical scenarios. Furthermore, current DL models fail to generalize to multi-class classification on varied, high-resolution datasets [5].

Research Gap: Though many DL models are suggested for DR detection, they are mostly limited by catastrophic forgetting in ongoing learning, poor noise handling in retinal fundus images, and weak generalization across multi-class severity levels. Not many experiments integrate successful preprocessing, stable feature extraction, and ongoing learning mechanisms within a single framework for credible DR detection.

Research Objectives: The proposed research seeks to engineer a strong and scalable DR detection system that (i) enhances image quality with high-end denoising methods, (ii) derives feature-rich spatial information for precise multi-class classification, and (iii) employs continual learning techniques to address catastrophic forgetting.

Significance of the Study: By overcoming the limitations of existing models, the new model can assist ophthalmologists with faster, accurate, and automated DR diagnosis. This will eventually lead to timely treatment, less vision loss, and better patient outcomes in diabetes care.

The key contributions of the proposed framework are depicted below:

- ❖ Enhance DR detection using a Color Wiener Filter (CWF) for denoising and a novel LPTCN for capturing detailed spatial features.
- ❖ Ensure continual learning by integrating EWC and HSR to prevent catastrophic forgetting.
- ❖ Validate the framework on the IDRID dataset using metrics like Accuracy, MCC, YI, SMAPE, and Computation Time, showing its superiority over existing methods.

The remaining part of the manuscript is prearranged as follows: Section 2 reviews related work, Section 3 describes the proposed method, Section 4 emphasizes results and discussion, and Section 5 concludes the study.

2. RELATED WORK: A BRIEF REVIEW

Several studies have investigated automated detection and classification of DR using fundus retinal images leveraging deep learning (DL), optimization algorithms, and attention mechanisms.

Deep learning architecture for DR detection

Jabbar *et al.* [6] suggested an efficient service platform for DR diagnosis, wherein RFI were analyzed using a computing unit to calculate the severity of the disease. In like manner, Jabbar *et al.* [7] used GoogleNet and ResNet architectures with APSO for their optimization, exhibiting enhanced feature extraction and classification performance. Gupta *et al.* [8] subsequently improved DL-based classification using the Fennec Fox Optimizer, creating the FIDRC-DLFFO model for AI-driven DR severity classification. Although these models attained competitive accuracy, their reliance on enormous computational resources and large-scale annotated data restricts scalability.

Attention-based and Hybrid Models

Romero-Oraá *et al.* [9] presented an attention mechanism that treated illuminated and shadowed retinal structures separately after a preprocessing procedure in order to better distinguish fine lesions. Hemanth *et al.* [10] put forward a hybrid method by implementing He Weighted BiLSTM (HWBLSTM) along with transfer learning for detecting DR and lesion segmentation with the Enhanced Grasshopper Optimizer Region Growing Algorithm (EGORGA). Dimensionality reduction was applied with Modified Singular Value Decomposition (MSVD) and lesion classification was conducted with SqueezeNet. These methods enhanced detection accuracy but are plagued with high model complexity and prolonged training times, which obfuscate real-time usage in clinical pathways.

From these previous works, it is clear that although DL and optimization-based models are robust in classification performance, they still struggle with certain critical issues. Noise and variability in RFI are still inadequately handled, lowering robustness. Multi-class classification at various stages of DR severity is still challenging. All techniques do not have mechanisms to cope with continual learning, resulting in catastrophic forgetting on introducing new data.

Positioning of the Proposed Work and Research Gap:

To address these challenges, the suggested framework incorporates a Color Wiener Filter (CWF) for noise removal, a new LPTCN network for fine-spatial feature extraction, and Elastic Weight Consolidation (EWC) and Hard Sample

Replay (HSR) to avoid catastrophic forgetting in continual learning. Unlike other models, this integrated method is developed to provide robust, scalable, and clinically usable DR detection at multiple levels of severity.

3. PROPOSED METHODOLOGY

This manuscript presents a novel continual DL framework for efficient DR disease detection. Initially, input retinal fundus images are taken from the IDRID Dataset for accurate DR disease detection which undergoes the preprocessing

stage by employing a CWF that can enhance image clarity by adaptively removing noise while maintaining edge details for further processing. After preprocessing, a novel LPTCN is introduced which effectively learns intricate spatial relationships within fundus images, leading to a more accurate distinction between different DR abnormalities such as non-DR, mild NPDR, moderate NPDR, severe NPDR, and PDR. To mitigate this issue, the proposed framework integrates EWC and HSR. The proposed framework's workflow is depicted in Figure 1.

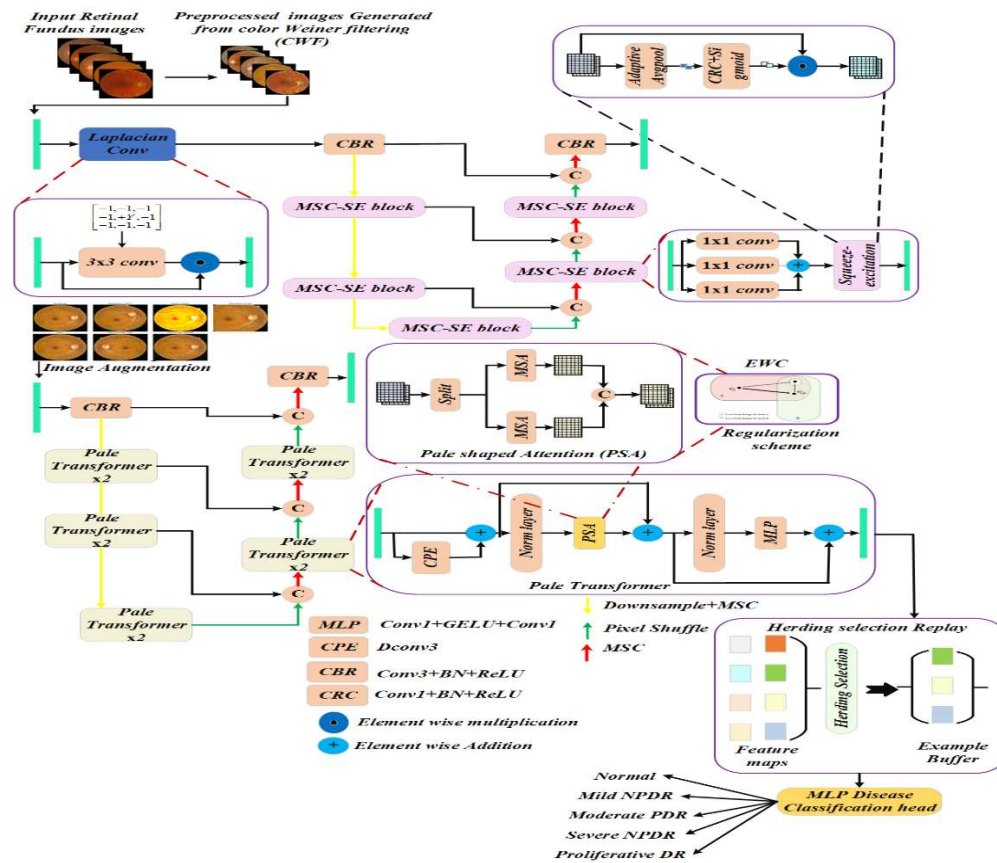


Figure 1: Workflow of the Proposed Framework

A. Data Acquisition

Initially, the IDRiD dataset [11] collected from freely accessible sources contains 516 high-resolution fundus images (4288×2848 pixels) captured in an Indian eye clinic. It includes segmentation (81 images with lesion masks), disease grading (516 images labeled for DR and DME severity), and localization (optic disc and fovea positions).

B. Preprocessing Stage

The acquired images from the database consist of high noise that cannot be directly fed into the proposed network model. To manifest this issue, a color wiener filtering (CWF) technique is introduced that preserves useful information and improves the illumination level of the neuroimages. In the RGB color space, the relationship between the actual image, the recovered image, and the additive noise can be expressed mathematically as follows:

$$x_i = Z_i + \gamma_i \quad (1)$$

Here, x_i indicates the observed instance which deliberates the pixel vector in the actual image Z_i , γ_i signifies the additive noise, and i denotes the image pixel which can be depicted as, $i(a, b)$ whereas a and b manipulates the pixel directions. $\tilde{x}$ manipulates the mean vector of the observed sample x_i . $\tilde{Z}$ characterizes the adjusted image of the pixel vector, and $\bar{\tilde{Z}}$ embodies the adjusted image of the mean vector which can be formulated using CWF $\mathfrak{R}$ is depicted below:

$$\bar{\tilde{Z}} = \hat{Z} = \mathfrak{R}(x_i - \tilde{x}) \quad (2)$$

Here, $\mathfrak{R}$ is evaluated to determine the mean square error (MSE) amid real and adjusted samples whereas MSE is mathematically articulated as,

$$MSE = \min \left| \left(\hat{Z} - \bar{\tilde{Z}} - \mathfrak{R}(x_i - \tilde{x}) \right)^2 \right| \quad (3)$$

In the actual image, when it fails to correlate with noise then the adjusted image $\mathfrak{R}$ can be formulated as,

$$\mathfrak{R} = m_{\text{cov}} (m_{\text{cov}} + m_{nn}) \quad (4)$$

Here, m_{nn} signifies the noise, and m_{cov} manipulates the covariance matrices. At last, the WF $\mathfrak{R}$ can be depicted mathematically as,

$$\mathfrak{R} = (m_{xx} - x_{nn}) m_{xx}^{-1} \quad (5)$$

Algorithm 1: Pseudocode for Proposed Framework

INPUT: IDRiD dataset fundus Images

$\{Z_1, Z_2, \dots, Z_n\}$

OUTPUT: DR severity classification

$\{\text{No_DR, Mild, Moderate, Severe, PDR}\}$

STEP 1: Data Acquisition

Load IDRiD dataset (516 fundus images, resolution 4288×2848)

STEP 2: Preprocessing using Color Wiener Filter (CWF)

FOR every image Z_i in Dataset **DO**

x_i = Observed pixel vector of Z_i

// Eq. (1)

γ_i = Additive noise

// Eq. (1)

Estimate adjusted pixel vector:

// Eq. (2)

$\tilde{Z} = \hat{Z} = \mathfrak{R}(x_i - \tilde{x})$

Calculate Mean Square Error (MSE)

// Eq. (3)

$MSE = \min \left| \left(\hat{Z} - \bar{\tilde{Z}} - \mathfrak{R}(x_i - \tilde{x}) \right)^2 \right|$

If Z_i uncorrelated with γ_i THEN

// Eq. (4)

$\mathfrak{R} = m_{\text{cov}} (m_{\text{cov}} + m_{nn})$

Final Wiener Filter output:

// Eq. (5)

$\mathfrak{R} = (m_{xx} - x_{nn}) m_{xx}^{-1}$

Store preprocessed image

END FOR

STEP 3: Feature Extraction with LPTCN

FOR each preprocessed image **DO**

Compute Laplacian operator for

edges: // Eq. (6)

$\nabla^2 g(u, v) = g(u+1, v) + g(u-1, v) + g(u, v+1) + g(u, v-1) - 4g(u, v)$

Extracted edge features with LCNN:

// Eq. (7)

$e = \tilde{r} (p \times \text{Laplacian_CNN}(\beta_0))$

Combine edge + intensity features

END FOR

STEP 4: Pale Transformer (PT) for Global Attention

FOR each Feature **DO**

Compute single-head self-attention:

// Eq. (8)

$V_r^a = \text{soft max} \left(\frac{\psi_x(U_c^a) \cdot \psi_k(U_c^a)^T}{\sqrt{D_k}} \right) \psi_v(U_c^a)$

Apply pale attention + position

coding: // Eq. (9–11)

$\hat{U}^l = U^{l-1} + \text{CPE}(U^{l-1})$

$\hat{U}^l = \hat{U}^l + \text{PA}(\text{LN}(\hat{U}^l))$

$U^l = \hat{U}^l + \text{MLP}(\text{LN}(\hat{U}^l))$

STEP 5: Continual Learning with HSR

FOR each class **DO**

Compute the class mean vector

```

    Select representative samples (RS):
    Store RS in the memory buffer for
    replay
  END FOR
STEP 6: Continual Learning with EWC
  FOR each task DO
    Compute Fisher Information Matrix
     $G$ 

    Regularized loss function:
    // Eq. (13)

     $\uparrow\downarrow(\phi) = \uparrow\downarrow_Y(\phi) + \sum_a \frac{\alpha}{2} G_a (\phi_a - \phi_{x,a}^*)$ 

  END FOR
STEP 7: Classification
  Train the final classifier on features
  Predict severity level  $\in \{\text{No\_DR}, \text{Mild}, \text{Moderate}, \text{Severe}, \text{PDR}\}$ 
RETURN Predicted DR Severity Levels

```

involves convolving the image with a 3×3 kernel.

$$\begin{aligned} \nabla^2 g(u,v) &= \frac{\partial^2 g}{\partial u^2} + \frac{\partial^2 g}{\partial v^2} \\ &= g(u+1,v) + g(u-1,v) + g(u,v+1) - 4g(u,v) \end{aligned} \quad (6)$$

Finally, a 3×3 convolution kernel incorporating the Laplacian operator is applied to determine gradient weights across eight different directions, facilitating the extraction of edge structures from DR images reconstructed via the FBP layer. The resulting edge information is then weighted by multiplying it by the original DR image. Due to the significant sparsity of high-frequency edge features, the MSC is combined with the SE mechanism to develop the multi-scale convolutional squeeze excitation (MSC-SE) block. The MSC block consists of triple convolutional (Conv) kernels of dimensions 1×1 , 3×3 and 5×5 , which integrate extracted feature information through elementwise accumulation. Furthermore, the SE block within the EEN component functions as a channel attention (CA) operation. It begins by applying an adaptive average pooling layer (AAPL) to reduce the longitudinal extents of the input attributes.

The outcome of the Laplacian CNN (LCNN) is depicted as,

$$e = \tilde{\gamma} \left(p \times \text{Laplacian_CNN}(\beta_0) \right) \quad (7)$$

Here, e indicates the outcome of edge features from the LCNN. The kernel conv weight within the LCNN module is constant and predefined throughout the learning phase. The LCNN is configured with the learnable parameter p for effective outcomes. Initially, its first two layers focus on capturing narrow pixel attributes using a pair of CBR blocks. Subsequently, after resolution reduction, the pixel attributes are resized to dimensions $N \times 128 \times 128 \times 128$ through two conv layers. At this stage, the LCNN module integrates the CBR unit with the PT unit to effectively seize global attribute representations.

b. Pale Transformer Module

Contrasting the conventional self-attention (SA) mechanism in ViT, the PT introduces a DConv-based adapted PS-Attention, where self-

C. DR Classification using LPTCN-EWC-HSR Technique

The preprocessed images are then fed into a continual learning (CL)-based Laplacian Pale Transformative Convolutional network (LPTCN), which integrates CL algorithms Elastic Weight Consolidation (EWC) and Herding Selection Replay (HSR) to prevent catastrophic forgetting. The architecture includes three main components: a data augmentation module, the LPTCN-EWC-HSR model, and a CL strategy. The augmentation module increases data diversity through transformations like rotation, cropping, blurring, and noise, with randomized execution to enhance variation. After normalization, the data is processed by the model, which features a backbone for extracting meaningful features and a classification layer for disease prediction based on learned patterns.

a. Primary Feature Representation Layer

For the extraction phase, the Laplacian Pale Transformative Convolutional network (LPTCN) is introduced. The novel Laplacian operator is emphasized, which helps to extract the edge-level features considering the second-order differential operator by determining the parallel and perpendicular gradient values of a sample positioned at the image's focal point region. Mathematically, the second-order gradient at each pixel location, represented as $g(u,v)$, is computed using the Laplacian operator, which

attention is computed within a pale-shaped region. This approach leverages DConv to lower computational complexity and reduce memory consumption. Additionally, self-attention is performed separately in token clusters arranged by rows and columns using stream segregation.

Then, the mapping of attention scores related to series and parallel features is generated using individual multi-head SA (MSA). The mathematical expression of single-head SA is depicted as,

$$V_r^a = \text{soft} \max \left(\frac{\psi_x(U_r^a) \cdot \psi_k(U_r^a)^T}{\sqrt{D_k}} \right) \psi_v(U_r^a), V_c^a$$

$$= \text{soft} \max \left(\frac{\psi_x(U_c^a) \cdot \psi_k(U_c^a)^T}{\sqrt{D_k}} \right) \psi_v(U_c^a)$$

(8)

Here, D_k indicates the normalization parameter for $\psi_x(U_r^a) \cdot \psi_k(U_r^a)^T$. At last, the evaluated SA attributes of series and parallel are combined with the channel extent to generate the outcome U . The entire PT block comprises triple series parts namely the pale attention (PA) block, the location coding block for obtaining the condition of location embedding, and the feature mapping using MLP. The process held in PT is mathematically articulated as,

$$\hat{U}^l = U^{l-1} + \text{CPE}(U^{l-1}) \quad (9)$$

$$\hat{U}^l = \hat{U}^l + \text{PA}(\text{LN}(\hat{U}^l)) \quad (10)$$

$$U^l = \hat{U}^l + \text{MLP}(\text{LN}(\hat{U}^l)) \quad (11)$$

Here, CPE indicates the 3×3 Dconv layer, LN represents the layer normalization, and the MLP comprises 1×1 conv layers compared to conventional fully connected (FC) layers. Finally, the features of the image are fused, which consists of CDR and 1×1 conv layers.

c. Herding Selection Replay (HSR)

The classic replay method helps prevent forgetting in continual learning by retaining key samples from past training. A variation, HSR, selects representative samples closest to each

class's feature space center. It up-samples features, computes class-wise mean vectors, and chooses the nearest samples to store in memory. This balanced approach ensures efficient learning and adapts as new classes are introduced.

d. Elastic Weight Consolidation (EWC)

EWC supports CL by limiting changes to synapses critical for earlier tasks. The parameters (weights and biases) for tasks ξ_X and ξ_Y are represented as X and Y . The optimal parameter sets that minimize errors for these tasks are also labeled X and Y , respectively. A solution can be found here X aligns with X and Y with Y .

Since computing the exact posterior is infeasible, EWC approximates it as a Gaussian distribution centered at the parameter X . The Fisher Information Matrix is utilized to calibrate the oblique accuracy. The corresponding loss function is defined in Equation (13),

$$\hat{\mathcal{J}}(\phi) = \hat{\mathcal{J}}_Y(\phi) + \sum_a \frac{\alpha}{2} G_a (\phi_a - \phi_{X,a}^*)^2$$

(13)

Here, $\hat{\mathcal{J}}_Y(\phi)$ interprets loss function for the task ξ_Y , $(.)$ manipulates the weight factor, $\phi_{X,a}^*$ manipulates the parameter after task learning, G contemplates the Fisher information matrix, α indicates the attributes that define the relation between new and old tasks.

4. RESEARCH METHOD AND EXECUTION PROTOCOL

The proposed scheme is analysed and processed via the Python simulation platform. For the investigation progression, hyperparameters like a learning rate of 0.001, min-batch size of 64, dropout of 0.2, and 100 epochs are considered. The proposed method is executed on Intel(R) Core (TM) i5-4300M CPU with 4GB installed RAM using a 64-bit operating system. For dataset utilization, 80% of the data is used for training, 10% for testing, and 10% for validation, maintaining the ratio of 8:1:1.

5. RESULTS AND DISCUSSION

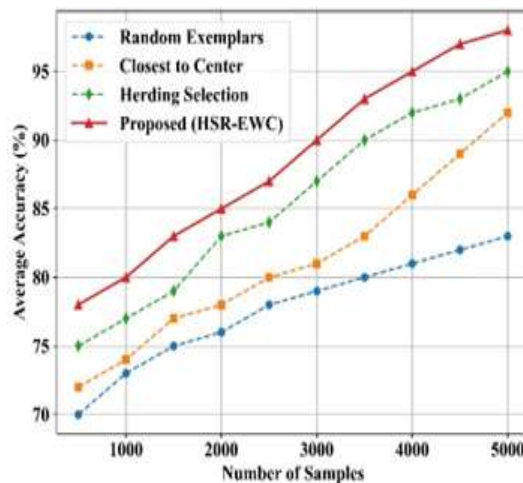
This section provides experimental results of the proposed model, such as confusion matrix assessment and comparative performance evaluation against models already implemented. A detailed discussion is offered to clarify the results achieved, identify strengths, and analyze the misclassifications or anomalies.

A. Confusion Matrix Analysis

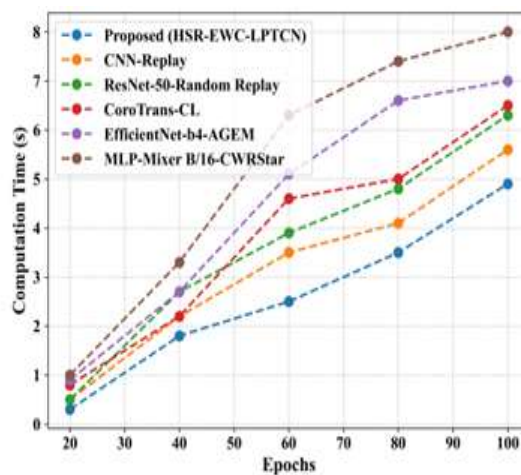
In this section, the Confusion matrix (CM) is analyzed for the developed framework under normal and diseased classes.

	MINPDR	MoNPDR	Normal	PDR	SNPDR
MINPDR	196	0	2	0	0
MoNPDR	0	196	0	1	1
Normal	1	0	196	1	0
PDR	1	1	0	196	0
SNPDR	0	0	1	1	196

Figure 2: CM Analysis under Varying Classes



(a)



(b)

Figure 3: Analysis of (a) Average Accuracy Analysis under Varying Samples, (b) Computation Time under Varying Epochs

Figures 3(a) and 3(b) compare sample selection accuracy and computation time. The proposed HSR-EWC method achieves the highest accuracy across all sample sizes and shows the most improvement as samples increase. It also

The confusion matrix (CM) in Figure 2 shows the performance of a classification model across five classes: Mild NPDR, Moderate NPDR, Normal, PDR, and SNPDR, with each class having 200 samples. The model correctly classified 196 samples in each class, indicating high accuracy. Only a few misclassifications occurred: Mild NPDR had 2 misclassified as Normal, Moderate NPDR had 1 as PDR and 1 as Severe NPDR, and Normal had 1 each misclassified as Mild NPDR and PDR. PDR had 1 misclassified as Mild NPDR and 1 as Moderate NPDR, while Severe NPDR had 1 misclassified as Normal and 1 as PDR.

B. Comparative Investigation of the Suggested Method over Other Schemes

In this section, the effectiveness achieved by the introduced method over the existing schemes is deliberated via graphical illustration. Several existing methods, like CNN-Replay, ResNet-50-Random Replay, CoroTrans-CL, EfficientNet-b4-AGEM, and MLP-Mixer B/16-CWRStar techniques, are compared with the proposed LPTCN-EWC-HSR technique framework. The comprehensive scrutiny of the attained efficacy is depicted below.

records the lowest computation time across epochs, highlighting its efficiency. In contrast, MLP-Mixer B/16-CWRStar and EfficientNet-b4-AGEM are slower and less accurate, especially with more epochs.

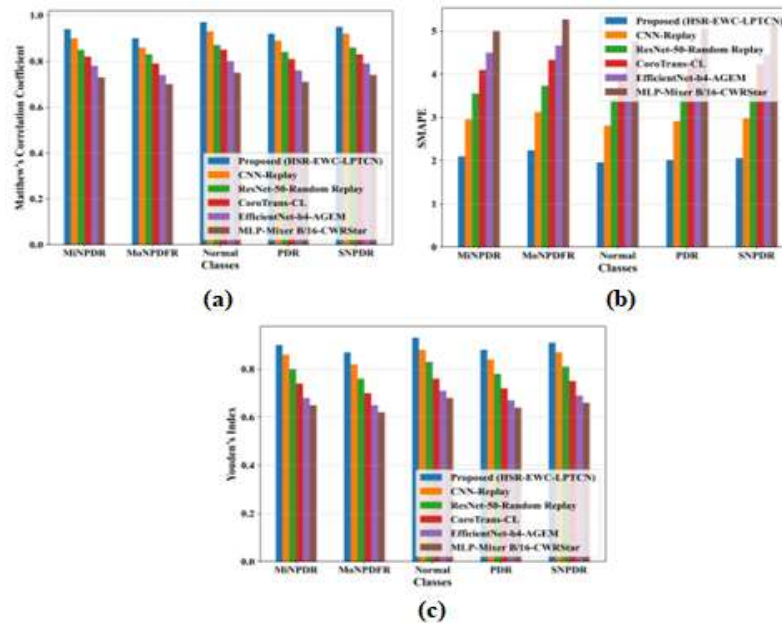


Figure 4: Analysis Of (A) MCC, (B) SMAPE, And (C) YI Under Varying Classes

Figures 4(a)–4(c) illustrate the performance of various models across MCC, SMAPE, and YI metrics. The proposed HSR-EWC-LPTCN method consistently outperforms all others, achieving the highest MCC and YI (above 0.85) and the lowest SMAPE (2.0–2.5) across all DR classes. CNN-Replay and ResNet-50-Random Replay show competitive results, while CoroTrans-CL and EfficientNet-b4-AGEM perform moderately. MLP-Mixer B/16-CWRStar ranks lowest, especially in complex classes like PDR and Severe NPDR. These results highlight the proposed model's robustness and classification accuracy.

C. Discussion

The comparison results in Figures 3 and 4 clearly show that the proposed HSR-EWC-LPTCN framework performs better consistently than the state-of-the-art methods, including CNN-Replay, ResNet-50-Random Replay, CoroTrans-CL, EfficientNet-b4-AGEM, and MLP-Mixer B/16-CWRStar. Various reasons account for this enhanced performance. State-of-the-art CNN-based models such as CNN-Replay and ResNet-50-Random Replay yield robust baseline performance but are poor at capturing fine retinal lesion boundaries, particularly in higher DR stages like PDR and Severe NPDR. The suggested model combines Laplacian Pyramid-based feature extraction (LPTCN), enhancing local texture and vascular structure

representation. It is thus able to detect subtle variations of lesions more effectively, leading to improved MCC and YI values. EfficientNet-b4-AGEM and MLP-Mixer B/16-CWRStar models experience catastrophic forgetting upon exposure to different classes of DR severity, resulting in poor performance in complicated categories. By implementing Elastic Weight Consolidation (EWC), the proposed method preserves essential knowledge throughout successive training epochs to avoid overfitting and guarantee stable classification despite increased sample sizes.

Figure 3(b) indicates that the proposed method consistently measures lower computation time per epoch than EfficientNet-b4-AGEM and MLP-Mixer B/16-CWRStar. This efficiency is because of hybrid sparse regularization (HSR), which removes unnecessary parameters and keeps valuable weights, thus speeding up convergence without loss of accuracy. Even though the confusion matrix indicates high classification accuracy (>98%), there are a few misclassifications. For instance, Mild NPDR cases sometimes overlap with Normal, while PDR cases are sometimes misclassified as Moderate NPDR. These mistakes can be related to visual similarities at early-stage lesions or ill-defined lesion progression, where boundaries are not clear-cut. It implies that, though the model is robust, further incorporation of clinical prior knowledge (e.g., ophthalmologist grading) could

be even more effective in reducing such mistakes. The systematically higher MCC (>0.85) and lower SMAPE (≈ 2.0 – 2.5) reflect the model's robustness in all DR classes, reflecting reliable generalization outside training data. In comparison to current models, the suggested framework not only boosts accuracy but also facilitates quicker execution, which is crucial for real-time DR screening in healthcare environments.

6. CONCLUSION

This work introduced a CDL-based framework for auto-Diabetic Retinopathy (DR) detection, incorporating Color Wiener Filter (CWF) preprocessing, the new LPTCN model for spatial feature extraction, and EWC-HSR mechanisms for dealing with catastrophic forgetting in continual learning. The proposed framework obtained better performance on the IDRID dataset with 97% average accuracy, 0.936 MCC, 2.07 SMAPE, 4.9s computation time, and 0.89 YI, outperforming other methods. These findings underscore the robustness, efficiency, and clinical potential of the proposed approach.

For future work, the model can be enhanced with explainable AI methods to improve interpretability, integration with multimodal clinical data to make diagnostic results more reliable, and real-time deployment in medical systems for large-scale screening. These directions will reinforce trust, usability, and clinical acceptability in real-world healthcare settings.

REFERENCE

- [1] M. D. Ramasamy, K. Periasamy, S. Periasamy, S. Muthusamy, P. Ramamoorthi, G. Thangavel, S. Sekaran, K. K. Sadasivuni, and M. Geetha, "A novel Adaptive Neural Network-Based Laplacian of Gaussian (AnLoG) classification algorithm for detecting diabetic retinopathy with colour retinal fundus images," *Neural Comput. Appl.*, vol. 36, no. 7, pp. 3513–3524, 2024.
- [2] F. Kallel and A. Echtioui, "Retinal fundus image classification for diabetic retinopathy using transfer learning technique," *Signal Image Video Process.*, vol. 18, no. 2, pp. 1143–1153, 2024.
- [3] A. Bilal, A. Imran, T. I. Baig, X. Liu, H. Long, A. Alzahrani, and M. Shafiq, "DeepSVDNet: A Deep Learning-Based Approach for Detecting and Classifying Vision-Threatening Diabetic Retinopathy in Retinal Fundus Images," *Comput. Syst. Sci. Eng.*, vol. 48, no. 2, 2024.
- [4] C. Nithyeswari and G. Karthikeyan, "An Effective Heuristic Optimizer with Deep Learning-assisted Diabetic Retinopathy Diagnosis on Retinal Fundus Images," *Eng. Technol. Appl. Sci. Res.*, vol. 14, no. 3, pp. 14308–14312, 2024.
- [5] G. Sivapriya, M. D. R. Manjula, P. Keerthika, and V. Praveen, "Automated diagnostic classification of diabetic retinopathy with microvascular structure of fundus images using deep learning method," *Biomed. Signal Process. Control*, vol. 88, p. 105616, 2024.
- [6] A. Jabbar, S. Naseem, J. Li, T. Mahmood, M. K. Jabbar, A. Rehman, and T. Saba, "Deep transfer learning-based automated diabetic retinopathy detection using retinal fundus images in remote areas," *Int. J. Comput. Intell. Syst.*, vol. 17, no. 1, p. 135, 2024.
- [7] A. Jabbar, H. B. Liaqat, A. Akram, M. U. Sana, I. D. Azpiroz, I. D. L. T. Diez, and I. Ashraf, "A lesion-based diabetic retinopathy detection through hybrid deep learning model," *IEEE Access*, 2024.
- [8] I. K. Gupta, S. Patil, S. Mahadevkar, K. Kotecha, A. K. Mishra, and J. JPC Rodrigues, "Retinal Fundus Imaging-Based Diabetic Retinopathy Classification using Transfer Learning and Fennec Fox Optimization," *MethodsX*, p. 103232, 2025.
- [9] R. Romero-Oraá, M. Herrero-Tudela, M. I. López, R. Hornero, and M. García, "Attention-based deep learning framework for automatic fundus image processing to aid in diabetic retinopathy grading," *Comput. Methods Programs Biomed.*, vol. 249, p. 108160, 2024.
- [10] S. V. Hemanth, S. Alagarsamy, and T. D. Rajkumar, "A novel deep learning model for diabetic retinopathy detection in retinal fundus images using pre-trained CNN and HWBLSTM," *J. Biomol. Struct. Dyn.*, pp. 1–19, 2024.
- [11] <https://www.kaggle.com/datasets/tinnkanjan-anuwat/idrid-dataset>