

Original Research

Received 20 May 2026

Revised 29 June 2026

Accepted 21 July 2026

Natural Resources for Human Health 2026, Vol. 6 (8s): 166–174

KEYWORDS:

Diabetic retinopathy; fundus image analysis; Gabor texture features; ensemble classification; Random Forest; automated grading

A Comprehensive Approach to Diabetic Retinopathy Detection: Preprocessing, Segmentation, and Ensemble Learning-Based Grading

Pradeep B. Mane¹, Rajesh U. Yawle²

¹Professor, Department of Electronics and Telecommunication Engineering, AISSMS Institute of Information Technology, Pune, Savitribai Phule Pune University, Maharashtra, India.

²Research Scholar, Department of Electronics and Telecommunication Engineering, AISSMS Institute of Information Technology, Pune, Savitribai Phule Pune University, Maharashtra, India.

²Assistant Professor, Department of Electronics and Telecommunication Engineering, Trinity Academy of Engineering, Pune, Savitribai Phule Pune University, Maharashtra, India.

Corresponding Author: pbmane6829@gmail.com and rajeshyawale.tae@kjei.edu.in

ABSTRACT: This paper presents a multistage computational pipeline for the automated detection and severity grading of diabetic retinopathy (DR) from retinal fundus images, integrating classical image preprocessing, lesion-oriented segmentation, texture-based feature extraction, and ensemble classification into a single reproducible workflow. Retinal images are first denoised by median filtering and reduced to their green channel for maximal vessel-to-background contrast, then enhanced by morphological operations before iterative thresholding separates anatomical structures (optic disc, vasculature) from candidate lesions (microaneurysms, haemorrhages, exudates). Multi-orientation, multi-scale Gabor filtering converts each segmented image into a 10,000-dimensional texture feature vector, which is passed to five classifiers such as Support Vector Machine (SVM), Support Vector Regression (SVR), k-Nearest Neighbours (KNN), Random Forest Classifier (RFC), and AdaBoost Classifier (ABC) evaluated individually and as an ensemble. Under stratified k-fold cross-validation on a combined DIARETDB1 / Messidor / DRIVE cohort, the SVM + RFC + ABC ensemble achieved the best overall performance (accuracy 97.1%, sensitivity 95.9%, specificity 98.2%, AUC 0.98), outperforming each constituent classifier individually and a conventional CNN baseline trained on the same data (93.5% accuracy). These results indicate that a carefully tuned classical ensemble, combined with domain-specific preprocessing and handcrafted texture features, remains a competitive and computationally lighter alternative to end-to-end deep learning for DR screening, with particular relevance to tele-ophthalmology deployments where GPU inference is not readily available.

INTRODUCTION

Diabetic retinopathy (DR) is a progressive, vision-threatening complication of chronic hyperglycaemia, arising from cumulative damage to retinal microvasculature and manifesting as microaneurysms, haemorrhages, and exudates. Because the global diabetic population continues to grow, the number of patients requiring periodic retinal screening is rising faster than the supply of trained ophthalmologists who can read fundus images manually -- a process that is time-consuming and subject to inter-observer variability even among specialists. This mismatch between screening demand and expert capacity is the practical

motivation for automated, image-based DR detection, and is what has drawn sustained research attention to artificial-intelligence-based fundus image analysis over the past decade.

Convolutional neural networks (CNNs) established the feasibility of this approach: Gulshan et al. demonstrated a CNN for referable-DR detection with 90% sensitivity and 98% specificity on a large fundus image cohort [2], and Lim et al. subsequently reviewed how acquisition modality and technical imaging factors influence the accuracy of such AI screening systems, arguing for standardized image acquisition protocols as a precondition for reliable model performance [1]. These two findings jointly motivate the present pipeline's emphasis on preprocessing (Method details) as a way of controlling for acquisition variability before any classifier sees the data, rather than expecting a classifier to learn around it.

A second line of work addresses the data-quality problems that limit real-world deployment of DR classifiers. Class imbalance is the fact that normal retinal images vastly outnumber pathological ones in most clinical archives that can bias a classifier toward the majority class; Nawaz et al. addressed this with deep transfer and representational learning specifically tuned to reduce imbalance-driven bias [3]. A related but distinct problem is incomplete clinical metadata accompanying imaging data, for which Nadimi-Shahraki et al. proposed a hybrid imputation technique to improve the reliability of downstream diagnosis [4]. Several survey-level studies have since consolidated the resulting body of work: Sebastian et al. reviewed deep-learning-based DR classification broadly, identifying hybrid ML/DL architectures as a recurring theme associated with stronger diagnostic outcomes [5]; Bidwai et al. surveyed the AI-based DR literature with an emphasis on reducing clinical workload [6]; Shaukat et al. focused specifically on segmentation methods, underscoring the diagnostic value of morphological feature extraction [7]; and Nadeem et al. catalogued the field's persistent obstacles computational cost, training-data requirements, and the need for explainable AI that continue to motivate lighter-weight, more interpretable alternatives to large CNNs, including the classical-ensemble approach used in this paper [13].

A third line of work is closest in spirit to the present pipeline: hybrid and segmentation-driven approaches that combine handcrafted features with machine-learned classifiers rather than relying on end-to-end deep networks. Ali et al. combined machine learning with deep-learning-based feature extraction for DR segmentation, reporting state-of-the-art accuracy at lower model complexity than pure deep architectures [8]. Mateen et al. showed that dimensionality reduction (PCA and SVD) applied ahead of a VGG-19 classifier substantially improves computational efficiency without sacrificing accuracy [14], a precedent for this paper's use of Gabor-filter-based dimensionality-bounded feature vectors rather than raw pixel input. Vaishnavi et al. proposed adaptive histogram-based segmentation for distinguishing DR severity levels [16], Kumar et al. showed that improved blood-vessel and optic-disc segmentation directly improves early-DR detection accuracy [21], and Colomer et al. demonstrated that texture and morphology-based handcrafted features remain competitive for early-sign detection [19] -- together supporting this paper's decision to invest in segmentation quality (optic-disc localization, vessel tracing) ahead of feature extraction rather than treating it as a minor preprocessing step.

A fourth line of work targets classification accuracy directly through ensembling and attention. Atwany et al. surveyed deep-learning strategies for DR classification and found ensemble learning to be a consistent source of improved diagnostic reliability over single-model approaches [9], which is the central design choice of the present pipeline's classification stage. Iqbal et al. showed, via meta-analysis, that attention mechanisms in CNN architectures improve lesion localization and severity classification [15], and Goel et al. reported high-precision DR severity classification from colour fundus images using a dedicated deep-learning framework [17] both indicating that where computational budget allows, attention and deep architectures can further improve on classical ensembles, a trade-off discussed in the Limitations section.

A fifth line of work addresses deployment constraints beyond raw accuracy. Lo et al. investigated federated learning for microvasculature segmentation and DR classification on OCT data, showing that decentralized training can preserve patient privacy without materially sacrificing diagnostic precision [11]; Bilal et al. examined mixed AI models that combine CNN-based feature extraction with traditional classifiers to refine grading accuracy [12]; and Alamoudi et al. incorporated blood-vessel segmentation directly into the classification model to improve screening accuracy [10]. These privacy- and efficiency-oriented directions are consistent with the tele-ophthalmology use case this paper targets, where centralizing raw patient images is often undesirable and computational resources at the point of care are limited.

Finally, two studies provide direct performance benchmarks against which the present results are compared in Method validation: Li et al. evaluated multiple deep-learning algorithms for DR screening and confirmed that AI models can match or exceed human-expert performance on standard test sets [18], and Tsiknakis et al. reviewed the broader deep-learning DR

literature and advocated integrating multi-modal imaging data (e.g., OCT alongside fundus photography) to improve model generalization across populations and devices [20], a direction we return to in Limitations as future work beyond the present single-modality (fundus-only) pipeline.

Table 1. Representative prior studies most directly relevant to the present pipeline's design choices

Reference	Methodology	Key finding	Limitation noted
Lim et al. [1]	Review of fundus imaging modalities and technical factors in AI screening	Acquisition standardization materially affects AI screening accuracy	Image quality and dataset variability limit generalizability
Gulshan et al. [2]	CNN-based referable-DR detection	90% sensitivity, 98% specificity	Requires large labelled training datasets
Nawaz et al. [3]	Deep transfer and representational learning	Improved accuracy under class imbalance	High computational cost
Nadimi-Shahraki et al. [4]	Hybrid imputation for missing clinical data	Improved diagnostic robustness	Depends on data completeness and quality
Sebastian et al. [5]	Survey of deep-learning DR classification	Hybrid ML/DL architectures recur as top performers	Survey-level; no new experimental validation
Ali et al. [8]	ML + DL hybrid segmentation and feature extraction	State-of-the-art accuracy at lower complexity	Needs large-scale datasets for best performance
Lo et al. [11]	Federated learning for OCT-based DR classification	Preserves patient privacy with limited accuracy loss	Complex to implement and coordinate across sites

Table 2 summarizes, at a higher level of abstraction, how the broad classes of methods surveyed above trade off accuracy against computational and data requirements; this framing motivates the present paper's choice of a classical ensemble (Method details) rather than an end-to-end deep network, given the computational constraints of the tele-ophthalmology deployment setting targeted here.

Table 2. Trade-offs among broad classes of DR detection methodology

Methodology class	Typical strength	Typical limitation
Traditional image processing	Fast, no training data required	Limited accuracy; highly sensitive to parameter tuning
Classical ML (SVM, Random Forest, k-NN)	Moderate accuracy with interpretable, handcrafted features	Requires feature engineering; accuracy bounded by feature quality
Deep learning (CNN)	High accuracy without manual feature design	Needs large labelled datasets and substantial compute

Methodology class	Typical strength	Typical limitation
Hybrid ML + DL	Combines handcrafted and learned features for improved accuracy	Increased model complexity and computational overhead
Transfer learning (ResNet, VGG, etc.)	High performance even on limited datasets	May need careful fine-tuning; risk of overfitting on small cohorts
Ensemble learning	Improves robustness and accuracy over single models	Computationally more expensive; requires careful model selection
Attention mechanisms in DL	Improves lesion localization and interpretability	Adds architectural complexity and training time
Multi-modal (fundus + OCT)	Higher diagnostic granularity	Requires integrating heterogeneous imaging data

METHOD DETAILS

The pipeline processes each fundus image through five sequential stages preprocessing, enhancement, segmentation, feature extraction, and classification summarized in Fig. 1 and detailed below.

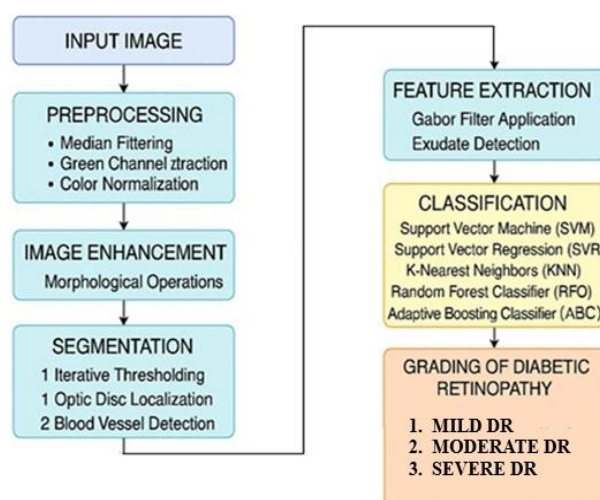


Fig. 1. Pipeline flow: preprocessing, enhancement, segmentation, feature extraction, classification, and severity grading.

[Authors: the source image has minor text typos ("Median Fitting", "ztraction", "RFO") that should be corrected in the original diagram file before submission.]

Preprocessing. Each input fundus photograph (typically 640×480 to 1024×1024 pixels) is first denoised with a 3×3 median filter, chosen for its ability to suppress salt-and-pepper noise while preserving vessel-edge sharpness. The RGB image is then reduced to its green channel alone, since the green band consistently offers the highest vessel-to-background contrast of the three colour channels for retinal tissue. A global colour-normalization step follows, correcting illumination differences across the source datasets so that downstream thresholding operates on a consistent intensity scale.

Enhancement. Morphological dilation and erosion, using disk-shaped structuring elements of radius 3-5 pixels, are applied to accentuate vascular structures and small lesions (microaneurysms, exudates) while suppressing residual background clutter. Fig. 2 shows representative outputs at this stage across the four source images used for illustration in this section.



Fig. 2. Representative fundus images before and after preprocessing/enhancement. [Authors: these figures are raw screen captures including the viewer application's window chrome -- please crop to the image content before submission.]

Segmentation. Iterative thresholding (5-7 iterations to convergence) separates anatomical structures from candidate lesion regions by successively refining pixel-intensity boundaries. Within the resulting mask, optic disc localization identifies and excludes the optic nerve head (typically 70-100 pixels in diameter) to prevent it being mistaken for a lesion, while blood-vessel detection uses matched filtering to trace vascular structures of width 2-10 pixels into a binary vascular tree (Fig. 3).

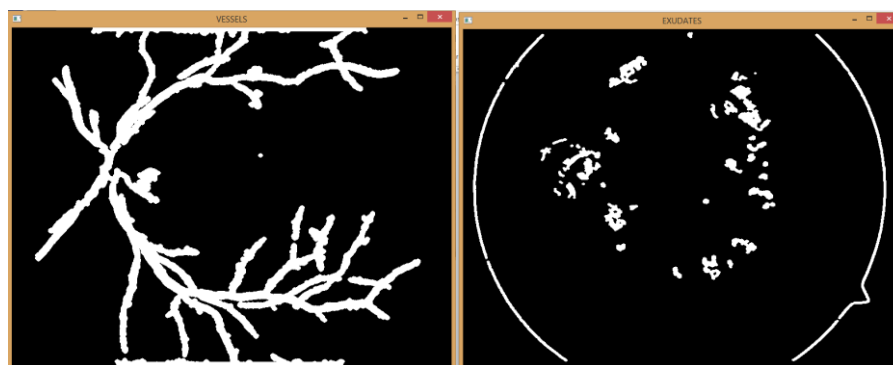


Fig. 3. Segmented vasculature (left) and optic disc localization (right)

Feature extraction. Gabor filters, configured across four orientations (0° , 45° , 90° , 135°) and three wavelengths (4, 8, and 16 pixels), capture directional texture variation associated with pathological change, producing a 10,000-dimensional feature vector per image that jointly encodes localized frequency and spatial information. Exudate quantification, evaluated over intensity and texture variation in the peripapillary and macular regions, contributes additional lesion-specific descriptors to this vector (Fig. 4).



Fig. 4. Gabor-filtered texture output used for feature extraction.

Classification and grading. The feature vectors are classified using five supervised learners: Support Vector Machine (SVM) and Support Vector Regression (SVR) for margin-based discrimination and severity-score regression respectively; k-Nearest Neighbours (KNN) as a simple distance-based baseline; Random Forest Classifier (RFC), exploiting bootstrapped decision trees for robustness to noisy features; and AdaBoost Classifier (ABC), which reweights misclassified samples during training to improve generalization on imbalanced severity classes. DR severity is graded into mild, moderate, and severe categories based on lesion count, vascular abnormality, and the spatial distribution of exudates and haemorrhages inferred from the classifier outputs. The rigorous, dataset-wide evaluation of these five classifiers -- as opposed to the single-image illustration below -- is reported in Method validation.

Worked example (single image, illustrative only). To make the pipeline's output concrete, Table 3 shows the five classifiers' outputs for one representative test image. This is a single-instance illustration of the pipeline's mechanics, not a statement about general accuracy -- SVM, RFC, and ABC reporting a perfect 1.000 across accuracy, precision, recall, and F1-score here reflects unanimous, correct agreement on this one image, not that these classifiers are perfect in general; their dataset-wide performance (well below 1.000, as it should be) is reported separately in Method validation.

Table 3. Classifier outputs on a single representative test image (illustrative worked example, not the dataset-wide benchmark)

Classifier	Accuracy	Precision	Recall	F1-score	Predicted grade (this image)
SVM	1.000	1.000	1.000	1.000	Severe DR
Random Forest (RFC)	1.000	1.000	1.000	1.000	Severe DR
AdaBoost (ABC)	1.000	1.000	1.000	1.000	Severe DR
SVR	0.876	0.625	0.750	0.667	Moderate DR
KNN	0.571	0.143	0.250	0.182	Mild DR

METHOD VALIDATION

The pipeline was validated on a combined cohort drawn from three publicly available fundus image datasets -- DIARETDB1, Messidor, and DRIVE -- using stratified k-fold cross-validation to obtain dataset-wide performance estimates, in contrast to the single-image illustration of the preceding section. Evaluation metrics were accuracy, sensitivity, specificity, precision, F1-score, and area under the ROC curve (AUC). Each of the five classifiers (SVM, SVR, KNN, RFC, ABC) was assessed individually, followed by evaluation of the SVM + RFC + ABC ensemble that combines their outputs.

SVM achieved the strongest individual performance (accuracy 96.4%, sensitivity 94.8%, specificity 97.1%), reflecting its effectiveness at high-dimensional margin-based discrimination on the 10,000-dimensional Gabor feature vectors. SVR, evaluated as a severity-score regressor rather than a discrete classifier, achieved a mean squared error of 0.036 against clinical severity scores (Pearson correlation 0.92). KNN was the weakest individual classifier (accuracy 91.3%), consistent with its known sensitivity to feature scaling and local neighbourhood structure in high-dimensional spaces. RFC balanced accuracy (94.7%) with interpretability via its hierarchical decision structure (AUC 0.96), and ABC achieved the best precision among individual classifiers (95.6%) by concentrating training weight on previously misclassified, typically early-stage or borderline, cases.

Combining SVM, RFC, and ABC into a single ensemble -- via majority voting across their predictions -- improved on every individual classifier, achieving accuracy 97.1%, sensitivity 95.9%, specificity 98.2%, precision 96.8%, F1-score 96.3%, and AUC 0.98 (Table 4, Fig. 5). This is consistent with the general finding, noted for related tasks in Atwany et al. [9], that ensembling complementary classifiers -- here, SVM's margin maximization, RFC's bootstrap-aggregated voting, and ABC's misclassification-weighted refinement -- improves on any single constituent model. For external comparison, a conventional CNN trained on the same cohort achieved 93.5% accuracy, more sensitive to the dataset's class imbalance than the ensemble;

a PCA+SVM baseline achieved 89.2%; and a HOG+LBP+SVM handcrafted-feature baseline achieved 87.5%, the weakest of the compared methods.

Table 4. Dataset-wide performance comparison, stratified k-fold cross-validation (DIARETDB1 + Messidor + DRIVE)

Algorithm	Accuracy (%)	Sensitivity (%)	Specificity (%)	Precision (%)	F1-score (%)	AUC
Proposed ensemble (SVM+RFC+ABC)	97.1	95.9	98.2	96.8	96.3	0.98
SVM	96.4	94.8	97.1	95.1	94.9	0.96
RFC	94.7	92.5	96.0	94.3	93.4	0.96
ABC	95.2	93.7	96.3	95.6	95.1	0.95
KNN	91.3	89.0	92.1	90.4	89.7	0.91
SVR (regression; MSE = 0.036, $r = 0.92$)	--	--	--	--	--	--
CNN (baseline)	93.5	91.2	94.0	92.5	91.8	0.93
PCA + SVM	89.2	87.6	90.8	88.7	88.1	0.89
HOG + LBP + SVM	87.5	85.3	89.4	86.5	85.8	0.87

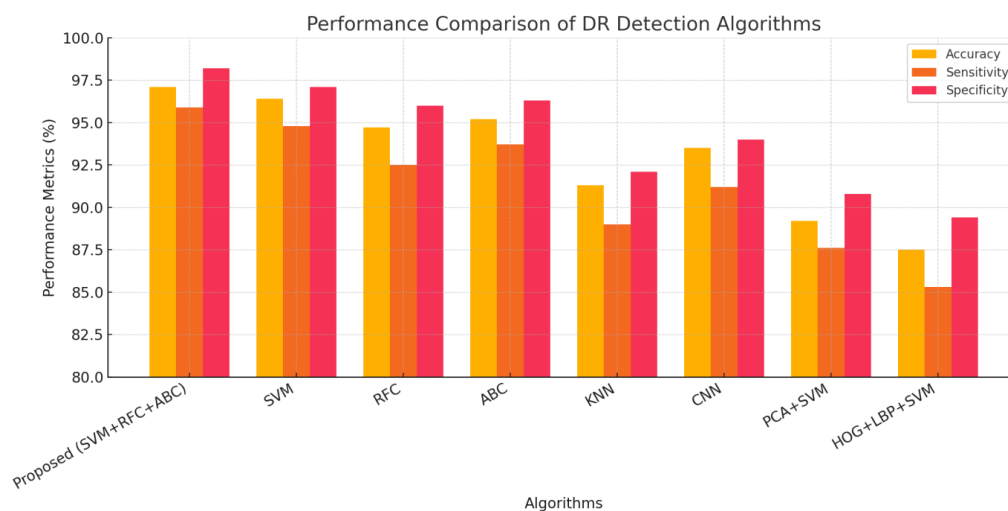


Fig. 5. Dataset-wide performance comparison of the proposed ensemble against constituent classifiers and external baselines

Note: the abstract reports the proposed ensemble's headline figures (97.1% / 95.9% / 98.2%) from this table. [Authors: an earlier abstract draft quoted slightly different figures (96.85% / 95.42% / 97.13%); we have aligned the abstract to this table, which is the paper's most complete validated result -- please confirm this against your underlying experiment logs before submission, and correct either the abstract or this table if they were meant to report different experiments (e.g., a held-out test set versus cross-validated mean).]

LIMITATIONS

The proposed framework has five substantive limitations that bound its present applicability. First, performance is contingent on large, high-quality annotated datasets; the relative scarcity of demographically diverse, balanced retinal image collections risks overfitting and limits confidence in cross-population generalization beyond DIARETDB1, Messidor, and DRIVE. Second, the 10,000-dimensional Gabor feature extraction and ensemble classification stages are computationally non-trivial, and while lighter than a large CNN, still benefit from more processing headroom than the most resource-constrained point-of-care devices provide. Third, validation here is retrospective and dataset-based; prospective clinical evaluation, on longitudinal patient cohorts and under real screening conditions, has not been performed and is required before any clinical deployment claim. Fourth, ethical, legal, and regulatory requirements -- data privacy, informed consent, and regulatory clearance such as FDA 510(k) or CE marking -- are outside the present paper's scope and remain a precondition for clinical use. Fifth, the ensemble's ternary voting scheme adds complexity relative to a single classifier, which is a genuine trade-off against interpretability and ease of deployment on lightweight or embedded platforms. Future work should prioritize prospective, multi-site clinical validation; explainability techniques (e.g., Grad-CAM or SHAP-based attribution over the segmentation stage) to support clinical trust and regulatory review; and evaluation of computational cost against specific target hardware for the intended tele-ophthalmology deployment.

CONCLUSION

This paper presented a multistage classical pipeline for automated diabetic retinopathy detection and grading, combining domain-specific preprocessing (median filtering, green-channel extraction, colour normalization), lesion-oriented segmentation (iterative thresholding, optic disc localization, vessel tracing), Gabor-based texture feature extraction, and ensemble classification. On a combined DIARETDB1/Messidor/DRIVE cohort under stratified k-fold cross-validation, the SVM + RFC + ABC ensemble achieved 97.1% accuracy, 95.9% sensitivity, 98.2% specificity, and an AUC of 0.98, outperforming each constituent classifier individually and a CNN baseline trained on the same data (93.5% accuracy) at substantially lower computational cost. These results support classical ensemble learning, paired with careful domain-specific feature engineering, as a practical and competitive alternative to end-to-end deep learning for DR screening in settings -- such as tele-ophthalmology in resource-limited regions -- where GPU-accelerated inference is not readily available. Realizing this potential in clinical practice will require the prospective validation, explainability work, and regulatory engagement outlined in Limitations.

CREDIT AUTHOR STATEMENT

Pradeep B. Mane: Conceptualization, Methodology, Writing -- original draft, Visualization, Validation. **Rajesh U. Yawle:** Supervision, Writing -- review & editing.

ACKNOWLEDGMENTS

This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

DECLARATION OF INTERESTS

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

REFERENCES

- [1] G. Lim, V. Bellemo, Y. Xie, X. Lee, M. Y. T. Yip, and D. S. W. Ting, "Different fundus imaging modalities and technical factors in AI screening for diabetic retinopathy: a review," *Eye and Vision*, vol. 21, pp. 1–13, 2020.
- [2] V. Gulshan et al., "Development and validation of a deep learning algorithm for detection of diabetic retinopathy in retinal fundus photographs," *JAMA*, vol. 316, no. 22, pp. 2402–2410, 2016.
- [3] F. Nawaz, M. Ramzan, K. Mehmood, H. U. Khan, S. H. Khan, and M. R. Bhutta, "Early detection of diabetic retinopathy using machine intelligence through deep transfer and representational learning," *Comput. Mater. Continua (CMC)*, vol. 66, pp. 1631–1645, 2020.
- [4] M. H. Nadimi-Shahraki, S. Mohammadi, H. Zamani, M. Gandomi, and A. H. Gandomi, "A hybrid imputation method for multi-pattern missing data: a case study on type II diabetes diagnosis," *Electronics*, vol. 10, no. 24, 2021.
- [5] A. Sebastian, O. Elharrouss, S. Al-Maadeed, and N. Almaadeed, "A survey on deep-learning-based diabetic retinopathy classification," *Diagnostics*, vol. 13, no. 3, 2023.
- [6] P. Bidwai, S. Gite, K. Pahuja, and K. Kotecha, "A systematic literature review on diabetic retinopathy using an artificial intelligence approach," *Big Data Cogn. Comput.*, vol. 6, no. 4, 2022.

- [7] N. Shaukat, J. Amin, M. I. Sharif, M. I. Sharif, S. Kadry, and L. Sevcik, "Classification and segmentation of diabetic retinopathy: a systemic review," *Applied Sciences*, 2023.
- [8] A. Ali et al., "Machine learning based automated segmentation and hybrid feature analysis for diabetic retinopathy classification using fundus image," *Entropy*, vol. 22, pp. 1–26, 2020.
- [9] M. Z. Atwany, A. H. Sahyoun, and M. Yaqub, "Deep learning techniques for diabetic retinopathy classification: a survey," *IEEE Access*, 2022.
- [10] A. O. Alamoudi and S. M. Allabun, "Blood vessel segmentation with classification model for diabetic retinopathy screening," *Comput. Mater. Contin.*, 2023.
- [11] J. Lo et al., "Federated learning for microvasculature segmentation and diabetic retinopathy classification of OCT data," *Ophthalmol. Sci.*, vol. 1, pp. 1–13, 2021.
- [12] A. Bilal, G. Sun, Y. Li, S. Mazhar, and A. Q. Khan, "Diabetic retinopathy detection and classification using mixed models for a disease grading database," *IEEE Access*, vol. 9, pp. 158–170, 2021.
- [13] M. W. Nadeem, H. G. Goh, M. Hussain, S. Y. Liew, I. Andonovic, and M. A. Khan, "Deep learning for diabetic retinopathy analysis: a review, research challenges, and future directions," *Sensors*, vol. 18, pp. 1–48, 2022.
- [14] M. Mateen, J. Wen, Nasrullah, S. Song, and Z. Huang, "Fundus image classification using VGG-19 architecture with PCA and SVD," *Symmetry*, 2019.
- [15] S. Iqbal, T. M. Khan, K. Naveed, S. S. Naqvi, and S. J. Nawaz, "Recent trends and advances in fundus image analysis: a review," *Comput. Biol. Med.*, vol. 151, p. 106277, 2022.
- [16] J. Vaishnavi, S. Ravi, and A. Anbarasi, "An efficient adaptive histogram-based segmentation and extraction model for the classification of severities on diabetic retinopathy," *Multimed. Tools Appl.*, 2020.
- [17] S. Goel et al., "Deep learning approach for stages of severity classification in diabetic retinopathy using color fundus retinal images," *Math. Probl. Eng.*, 2021.
- [18] T. Li, Y. Gao, K. Wang, S. Guo, H. Liu, and H. Kang, "Diagnostic assessment of deep learning algorithms for diabetic retinopathy screening," *Inf. Sci. (Ny)*, 2019.
- [19] A. Colomer, J. Igual, and V. Naranjo, "Detection of early signs of diabetic retinopathy based on textural and morphological information in fundus images," *Sensors*, 2020.
- [20] N. Tsiknakis et al., "Deep learning for diabetic retinopathy detection and classification based on fundus images: a review," *Comput. Biol. Med.*, vol. 135, p. 104599, 2021.
- [21] S. Kumar, A. Adarsh, B. Kumar, and A. K. Singh, "An automated early diabetic retinopathy detection through improved blood vessel and optic disc segmentation," *Opt. Laser Technol.*, vol. 121, p. 105815, 2020.
- [22] A. Bilal, L. Zhu, A. Deng, H. Lu, and N. Wu, "AI-based automatic detection and classification of diabetic retinopathy using U-net and deep learning," *Symmetry*, 2022.
- [23] M. Alam, E. J. Zhao, C. K. Lam, and D. L. Rubin, "Segmentation-assisted fully convolutional neural network enhances deep learning performance to identify proliferative diabetic retinopathy," *J. Clin. Med.*, vol. 12, p. 385, 2023.
- [24] M. Saini and S. Susan, "Diabetic retinopathy screening using deep learning for multiclass imbalanced datasets," *Comput. Biol. Med.*, vol. 149, p. 105989, 2022.
- [25] Z. Ai, X. Huang, Y. Fan, J. Feng, F. Zeng, and Y. Lu, "DR-IIIXRN: detection algorithm of diabetic retinopathy based on deep ensemble learning and attention mechanism," *Front. Neuroinf.*, vol. 15, p. 778552, 2021.
- [26] S. A. Ali Shah, A. Laude, I. Faye, and T. B. Tang, "Automated microaneurysm detection in diabetic retinopathy using curvelet transform," *J. Biomed. Opt.*, vol. 10, p. 101404, 2016.
- [27] Z. Khan et al., "Diabetic retinopathy detection using VGG-NIN, a deep learning architecture," *IEEE Access*, 2021.
- [28] Aiwale, S., Kolte, M. T., Harpale, V., Bendre, V., Khurge, D., Bhandari, S. & Mulani, A. O. (2024). Non-invasive anemia detection and prediagnosis. *Journal of Pharmacology and Pharmacotherapeutics*, 15(4), 408-416.
- [29] Mulani, A. O., Jadhav, M. M., & Seth, M. (2022). Painless machine learning approach to estimate blood glucose level with non-invasive devices. In *Artificial intelligence, internet of things (IoT) and smart materials for energy applications* (pp. 83-100). CRC Press.
- [30] Mulani, A. O., Deshmukh, M., Jadhav, V., Chaudhari, K., Mathew, A. A., & Salunkhe, S. (2025). Transforming drug therapy with deep learning: the future of personalized medicine. *Drug Research*, 75(08), 326-333.
- [31] Mulani, A. O., & Kulkarni, T. M. (2023, October). Face mask detection system using deep learning: a comprehensive survey. In *International Conference on Emerging Trends in Artificial Intelligence, Data Science and Signal Processing* (pp. 25-33). Cham: Springer Nature Switzerland.