

An Efficient Deep Learning Framework for Diabetic Retinopathy Classification using Generative Data Augmentation and Knowledge Distillation

by

TASNIMUL MOHAMMAD FAHIM

21301467

TANZIM HOSSAIN AYON

21201531

SHUVO DATTA

21101241

JAWAD BIN ALAM

20101040

SARIKA BINTEY SABUR

20101074

A thesis submitted to the Department of Computer Science and Engineering
in partial fulfillment of the requirements for the degree of
B.Sc. in Computer Science and Engineering

Department of Computer Science and Engineering
Brac University
October 2025

© 2025. Brac University
All rights reserved.

Declaration

It is hereby declared that

1. The thesis submitted is our own original work while completing degree at Brac University.
2. The thesis does not contain material previously published or written by a third party, except where this is appropriately cited through full and accurate referencing.
3. The thesis does not contain material which has been accepted, or submitted, for any other degree or diploma at a university or other institution.
4. We have acknowledged all main sources of help.

Student's Full Name & Signature:

Jawad Bin Alam

20101040

Sarika Bintey Sabur

20101074

Tanzim Hossain Ayon

21201531

Tasnimul Mohammad Fahim

21301467

Shuvo Datta

21101241

Approval

The thesis titled “An Efficient Deep Learning Framework for Diabetic Retinopathy Classification using Generative Data Augmentation and Knowledge Distillation” submitted by

1. Jawad Bin Alam (20101040)
2. Sarika Binte Sabur (20101074)
3. Tanzim Hossain Ayon (21201531)
4. Tasnimul Mohammad Fahim (21301467)
5. Shuvo Datta (21101241)

of Summer 2025 has been accepted as satisfactory in partial fulfillment of the requirement for the degree of B.Sc. in Computer Science and Engineering in Summer 2025.

Examining Committee:

Supervisor:
(Member)

Dr. Amitabha Chakrabarty

Professor
Department of Computer Science and Engineering
Brac University

Thesis Coordinator:
(Member)

Md. Golam Rabiul Alam, PhD

Professor
Department of Computer Science and Engineering
Brac University

Head of Department:
(Chair)

Sadia Hamid Kazi

Chairperson and Associate Professor
Department of Computer Science and Engineering
Brac University

Abstract

Diabetic retinopathy (DR) affects 93 million people globally and is a leading cause of preventable blindness. Automated screening systems face two critical barriers: severe class imbalance in medical datasets where healthy cases vastly outnumber vision-threatening stages, and high computational demands that limit deployment in resource-constrained settings. This research presents a three-stage framework integrating Progressive Growing GANs for synthetic fundus image generation, a dual-branch ensemble architecture combining EfficientNet-V2-S with Vision Transformer-B/16, and knowledge distillation to compress the model into MobileNetV2. The generative approach created a balanced training corpus of 10,000 images across five severity classes, addressing the fundamental imbalance problem while maintaining pathological authenticity validated through perceptual metrics. On APTOS 2019, the ensemble teacher achieves 95.8% accuracy with 0.975 quadratic weighted kappa. Cross-dataset validation on DDR and Messidor-2 confirms robust generalization across diverse clinical populations and imaging protocols. The distilled MobileNetV2 student retains 92.4% accuracy and 0.938 kappa while reducing model parameters by over ninety-six percent and computational operations by over ninety-seven percent. Inference latency decreases by over eighty percent on CPU, with model size compressed to enable mobile deployment. Grad-CAM visualizations confirm clinically relevant attention to microaneurysms, hemorrhages, and neovascularization. This framework demonstrates that generative augmentation combined with heterogeneous ensemble learning and knowledge distillation overcomes the accuracy-deployability trade-off, enabling accessible DR screening in resource-limited settings.

Keywords: Diabetic Retinopathy, Progressive Growing GAN, Knowledge Distillation, Vision Transformer, EfficientNet, Class Imbalance, Medical Image Analysis, Ensemble Learning

Table of Contents

Declaration	i
Approval	ii
Abstract	iv
Table of Contents	v
List of Figures	vii
List of Tables	viii
Nomenclature	viii
1 Introduction	1
1.1 Background	1
1.2 Motivation	2
1.3 Problem Statement	3
1.4 Objectives	3
1.4.1 Specific Objectives	3
1.5 Methodology Overview	4
1.5.1 Stage 1: Generative Data Augmentation	4
1.5.2 Stage 2: Dual-Branch Ensemble	4
1.5.3 Stage 3: Knowledge Distillation	4
1.5.4 Evaluation Protocol	4
1.6 Scope and Limitations	5
1.6.1 Scope	5
1.6.2 Limitations and Challenges	5
2 Literature Review	6
2.1 Preliminaries	6
2.1.1 Diabetic Retinopathy Classification	6
2.1.2 Neural Network Architectures	6
2.1.3 Generative Adversarial Networks	7
2.1.4 Knowledge Distillation	7
2.1.5 Evaluation Metrics	8
2.2 Literature Review	8
2.2.1 Deep Learning Architectures for DR Classification	8
2.2.2 Addressing Class Imbalance	10
2.2.3 Model Compression and Efficient Deployment	11

2.2.4	Cross-Dataset Generalization	11
2.2.5	Clinical Deployment Considerations	12
2.2.6	Summary of Key Findings and Research Gaps	13
3	Requirements, Impacts and Constraints	15
3.1	Final Specifications and Requirements	15
3.2	Societal Impact	16
3.3	Environmental Impact	16
3.4	Ethical Issues	17
3.5	Project Management Plan	17
3.6	Risk Management	17
3.7	Economic Analysis	18
4	Methodology	19
4.1	Dataset Description and Methodology	19
4.1.1	Dataset Description	19
4.1.2	Research Workflow	22
4.1.3	Image Preprocessing	22
4.1.4	Data Augmentation	22
4.1.5	Phase 1: Baseline Evaluation	23
4.1.6	Phase 2: Synthetic Data Generation	24
4.1.7	Phase 3: Comparative Analysis	25
4.1.8	Phase 4: Ensemble Construction	25
4.1.9	Phase 5: External Validation	26
4.1.10	Phase 6: Lightweight Model Evaluation	26
4.1.11	Phase 7: Knowledge Distillation	26
4.1.12	Phase 8: Comprehensive Evaluation	27
5	Result Analysis	28
5.1	Results and Analysis	28
5.1.1	Phase 1: Baseline Architecture Evaluation Results	28
5.1.2	Phase 2: Synthetic Data Quality Assessment	29
5.1.3	Phase 3: Performance on Balanced Dataset	29
5.1.4	Phase 4: Dual-Branch Ensemble Performance	32
5.1.5	Phase 5: External Validation on Messidor-2	33
5.1.6	Phase 6: Lightweight Model Evaluation	34
5.1.7	Phase 7: Knowledge Distillation Results	34
5.1.8	Phase 8: Comprehensive Final Evaluation	36
5.1.9	Computational Efficiency Analysis	37
5.1.10	Cross-Dataset Generalization Analysis	38
5.1.11	Statistical Reliability Analysis	38
5.1.12	Clinical Impact Assessment	38
6	Conclusion	40
6.1	Summary of Findings	40
6.2	Contributions to the Field	40
6.3	Recommendations for Future Work	41
	Bibliography	43

List of Figures

4.1	Class distribution showing severe imbalance in APTOS 2019 dataset.	20
4.2	DR progression samples (top: original; bottom: CLAHE-preprocessed).	21
4.3	Class distribution after GAN augmentation showing perfect balance.	21
4.4	Nine-phase research methodology from imbalanced dataset to deploy- able model.	22
4.5	Progressive Growing GAN architecture showing incremental resolu- tion growth.	24
4.6	Dual-branch ensemble with knowledge distillation to MobileNetV2.	26
4.7	Grad-CAM visualization: original image, heatmap, and overlay.	27
5.1	Training and validation loss curves for EfficientNet-V2-S on the bal- anced APTOS 2019 dataset.	30
5.2	Training and validation accuracy curves for EfficientNet-V2-S on the balanced APTOS 2019 dataset.	31
5.3	Training and validation loss curves for ViT-B/16 on the balanced APTOS 2019 dataset.	31
5.4	Training and validation accuracy curves for ViT-B/16 on the balanced APTOS 2019 dataset.	32
5.5	Training and validation loss curves for the dual-branch ensemble model.	33
5.6	Training and validation accuracy curves for the dual-branch ensemble model.	33
5.7	Training and validation loss curves for the distilled student model (MobileNetV2).	35
5.8	Training and validation accuracy curves for the distilled student model (MobileNetV2).	35
5.9	Validation loss comparison across all final models.	36
5.10	Validation accuracy comparison across all final models.	37

List of Tables

2.1	International Clinical Diabetic Retinopathy Disease Severity Scale (ICDRSS)	6
2.2	CNN architecture performance comparison	9
2.3	Vision Transformer performance	9
2.4	GAN-based augmentation in medical imaging	10
2.5	Knowledge distillation methods and compression	11
2.6	Cross-dataset validation results	12
4.1	Training Set (80%, 2,930 Images)	20
4.2	Validation Set (20%, 732 Images)	22
4.3	Final balanced training dataset composition.	22
4.4	Nine-Phase Research Workflow	23
4.5	Preprocessing Pipeline	23
4.6	Geometric Transformations	24
4.7	Photometric Adjustments	24
4.8	Advanced Regularization	25
5.1	Baseline CNN Architecture Performance on Imbalanced APTOS 2019	28
5.2	Baseline Vision Transformer Performance on Imbalanced APTOS 2019	29
5.3	Synthetic Image Quality Metrics by DR Class	29
5.4	Performance Comparison: Imbalanced vs. Balanced Training	30
5.5	Ensemble Model Performance and Fusion Analysis	32
5.6	Cross-Dataset Validation on Messidor-2 (1,748 images)	34
5.7	Lightweight Architecture Performance Comparison	34
5.8	Knowledge Distillation Performance with Varying Hyperparameters .	35
5.9	Final Model Performance Comparison on APTOS Validation Set . . .	36
5.10	Per-Class Performance Metrics for Ensemble Teacher Model	37
5.11	Computational Efficiency: Teacher vs. Student	38
5.12	Cross-Dataset Performance Without Fine-tuning	38

Chapter 1

Introduction

1.1 Background

Diabetic retinopathy (DR) represents the leading cause of preventable blindness among working-age adults globally, affecting approximately 93 million individuals [15]. The condition progresses through five severity stages defined by the International Clinical Diabetic Retinopathy Disease Severity Scale (ICDRSS): No DR, Mild nonproliferative DR (NPDR), Moderate NPDR, Severe NPDR, and Proliferative DR [3]. Each stage exhibits increasingly severe retinal abnormalities including microaneurysms, hemorrhages, exudates, and pathological neovascularization.

Early detection through systematic screening programs enables timely therapeutic interventions that preserve vision. However, traditional manual screening encounters significant barriers: critical shortage of trained ophthalmologists particularly in underserved regions, time-intensive examination processes requiring specialized equipment, and substantial inter-observer variability in severity grading [4]. These limitations prove especially acute in resource-constrained settings where DR prevalence peaks.

Deep learning has demonstrated human-level performance in automated medical image analysis. Convolutional neural networks (CNNs) including ResNet, DenseNet, and EfficientNet achieve robust DR classification through hierarchical feature learning [9]. The EfficientNet family models offer optimal accuracy-efficiency trade-offs through compound scaling that balances network depth, width, and resolution. Vision Transformers (ViT) leverage self-attention mechanisms to model long-range spatial dependencies, capturing global retinal patterns complementary to local lesion features extracted by CNNs [22].

Despite these advances, clinical translation faces two fundamental obstacles. First, severe class imbalance in public DR datasets mirrors real-world prevalence distributions: benchmark datasets like APTOS 2019 contain 49.3% No-DR cases while vision-threatening Severe DR represents only 5.3% of samples. This imbalance induces systematic classification bias toward majority classes, compromising detection of minority stages requiring urgent intervention [2]. Traditional resampling techniques prove inadequate for complex medical imaging where subtle pathological features distinguish severity levels.

Second, state-of-the-art models demand substantial computational resources unsuitable for primary care deployment where screening proves most beneficial. Large ensemble models and transformers, while delivering superior accuracy, require ex-

tensive memory and processing power that limit accessibility in resource-constrained environments. This creates a fundamental accuracy-efficiency trade-off preventing widespread adoption of AI-based screening systems.

Generative Adversarial Networks (GANs) offer solutions to class imbalance by synthesizing realistic minority-class samples [6]. Progressive Growing GANs (PG-GANs) stabilize high-resolution generation through gradual resolution scaling from 4×4 to 512×512 pixels, preserving fine pathological details essential for diagnosis [17]. Knowledge distillation compresses large models into efficient architectures by transferring probabilistic knowledge from complex teachers to lightweight students [7], substantially reducing computational requirements while maintaining diagnostic performance.

Recent DR research has explored transfer learning ensembles, architectural modifications, and compact transformers [25], [26]. However, no prior work systematically integrates class-conditional generative augmentation, heterogeneous CNN-Transformer ensemble learning, and knowledge distillation with comprehensive multi-dataset validation and explicit efficiency-accuracy trade-off analysis.

1.2 Motivation

The global DR burden continues escalating, with prevalence projected to reach 160 million by 2030. Critical workforce shortages compound this crisis: many regions lack sufficient ophthalmologists for screening demands. Automated AI-based systems offer transformative potential for improving DR screening accessibility, particularly in low-resource settings.

Current deep learning approaches face a fundamental paradox. Models achieving high accuracy require extensive computational resources and large balanced training datasets—both scarce in medical imaging. Severe class imbalance reflects real-world epidemiology but creates algorithmic bias: models excel at detecting common cases while failing to identify rare proliferative DR requiring urgent intervention. This undermines diagnostic utility and safety.

Existing solutions address these challenges in isolation. Generative augmentation lacks rigorous quality validation or integration with modern architectures. High-performance CNN-Transformer ensembles remain computationally prohibitive. Model compression through distillation is rarely evaluated on medical tasks requiring calibrated probability estimates for decision support.

This research addresses the need for a comprehensive framework simultaneously tackling class imbalance through validated generative synthesis, maximizing diagnostic accuracy through architectural synergy, and ensuring computational efficiency through intelligent compression. Such a system could enable DR screening in resource-constrained environments, potentially improving early detection of vision-threatening stages.

Furthermore, lack of standardized cross-dataset validation raises concerns about generalization. Many studies report performance on single datasets without testing external populations, limiting confidence in real-world applicability across diverse clinical settings, imaging protocols, and demographics. A robust framework must demonstrate consistent performance across multiple independent datasets to ensure reliability.

1.3 Problem Statement

Automated diabetic retinopathy classification systems face three interconnected challenges that limit their practical application:

Severe Class Imbalance: DR datasets exhibit extreme imbalance where "No DR" cases outnumber vision-threatening stages by ratios exceeding 9:1. This fundamentally compromises performance on minority classes—the advanced stages requiring urgent intervention. Traditional augmentation fails to capture complex pathological variations distinguishing severity levels, while simple oversampling creates unrealistic duplicates that do not improve model generalization.

Computational Inefficiency: State-of-the-art models achieving high accuracy require substantial resources (hundreds of millions of parameters, tens of GFLOPs per inference) impractical for mobile devices or resource-constrained environments. This contradicts practical reality: DR screening is most urgently needed in low-resource settings where high-performance computing infrastructure is unavailable. Current approaches optimize either accuracy or efficiency, not both simultaneously.

Limited Generalization Validation: Existing literature predominantly reports performance on single datasets without comprehensive cross-dataset validation. Models trained on specific geographic regions, imaging protocols, or demographics may not generalize to diverse real-world populations, raising critical questions about reliability and safety in heterogeneous healthcare settings.

These problems create a fundamental accuracy-efficiency trade-off that prevents widespread adoption of AI-based DR screening systems in the settings where they are most needed.

1.4 Objectives

The primary objective is to develop and validate a comprehensive framework for automated diabetic retinopathy classification that overcomes class imbalance and computational inefficiency while ensuring robust cross-population generalization.

1.4.1 Specific Objectives

1. Implement class-conditional Progressive Growing GAN to synthesize high-fidelity fundus images for minority DR classes, creating balanced training data. Validate synthetic quality using perceptual metrics (FID, LPIPS, Inception Score) and evaluate impact on minority class performance.
2. Design dual-branch ensemble combining EfficientNet-V2-S and ViT-B/16 to leverage local feature extraction and global context modeling. Optimize fusion strategy to maximize quadratic weighted kappa for ordinal classification.
3. Develop knowledge distillation framework compressing ensemble teacher into lightweight MobileNetV2 student. Evaluate accuracy-efficiency trade-offs measuring diagnostic metrics (accuracy, QWK, AUC) and system-level indicators (parameters, FLOPs, model size).
4. Assess generalization through zero-shot evaluation on external datasets (DDR, Messidor-2) representing diverse populations, geographic regions, and imaging

protocols without fine-tuning.

5. Implement Grad-CAM visualizations to verify model attention on clinically relevant pathological features and ensure interpretability.

1.5 Methodology Overview

The proposed framework comprises three integrated stages addressing class imbalance, diagnostic accuracy, and computational efficiency.

1.5.1 Stage 1: Generative Data Augmentation

Class-conditional Progressive Growing GAN is trained on APTOS 2019 to synthesize high-resolution fundus images for minority DR classes. The architecture progressively increases resolution from 4×4 to 512×512 pixels through seven stages, ensuring training stability while preserving pathological details essential for classification. Synthetic images create a balanced corpus of 2,000 images per class (10,000 total). Quality is validated using three perceptual metrics: Fréchet Inception Distance (FID) measuring distributional similarity, Learned Perceptual Image Patch Similarity (LPIPS) quantifying perceptual realism, and Inception Score assessing quality and diversity.

1.5.2 Stage 2: Dual-Branch Ensemble

EfficientNet-V2-S extracts local texture patterns through efficient convolutional operations, while Vision Transformer-B/16 captures global spatial relationships through self-attention mechanisms. Both branches utilize ImageNet pre-training followed by two-phase training: (1) freeze backbone and train classification head for 5-10 epochs; (2) end-to-end fine-tuning for 20-30 epochs. Training employs regularization techniques including Mixup augmentation and label smoothing. Predictions are fused using AUC-weighted averaging that automatically adapts to model reliability, outperforming simple averaging.

1.5.3 Stage 3: Knowledge Distillation

Ensemble knowledge is transferred to MobileNetV2 through temperature-scaled soft target matching. The distillation loss combines hard label cross-entropy with temperature-scaled KL divergence between teacher-student probability distributions. Optimal hyperparameters (temperature $T=3.0$, distillation weight $=0.2$) balance knowledge transfer and model compression, achieving 96.8% parameter reduction while retaining 93.9% of teacher accuracy.

1.5.4 Evaluation Protocol

Experiments utilize APTOS 2019 dataset with stratified 80% training and 20% validation split. Cross-dataset validation employs DDR (13,673 images from 147 Chinese hospitals) and Messidor-2 (1,748 images from French ophthalmologic departments) for zero-shot evaluation. Evaluation metrics include accuracy, quadratic

weighted kappa (QWK) for ordinal classification, AUC, F1-score, per-class precision and recall. Efficiency metrics quantify parameters (millions), FLOPs (giga-operations), and model size (MB). Grad-CAM visualizations provide interpretability by highlighting influential image regions.

1.6 Scope and Limitations

1.6.1 Scope

This research focuses on five-class DR severity grading from color fundus photographs using the ICDRSS classification system. The framework encompasses three primary components: class-conditional GAN synthesis for data balancing, dual-branch CNN-Transformer ensemble for classification, and knowledge distillation for model compression. Evaluation includes diagnostic performance metrics, computational efficiency analysis, cross-dataset generalization assessment, and interpretability through attention visualization.

1.6.2 Limitations and Challenges

Generative Quality Assessment: Synthetic image quality is validated using computational perceptual metrics (FID, LPIPS, Inception Score) rather than formal evaluation by multiple ophthalmologists. While these metrics provide objective measures of image quality and diversity, they may not fully capture all clinically relevant features.

Computational Resources: Training Progressive GANs and large ensemble models requires substantial GPU infrastructure (NVIDIA RTX 4090 with 24GB VRAM). This may limit reproducibility in resource-constrained research settings, though the final distilled student model requires minimal computational resources.

Dataset Generalizability: Cross-dataset validation on DDR and Messidor-2 provides strong generalization evidence across different populations, geographic regions, and imaging protocols. However, this may not fully represent global diversity in imaging equipment specifications, acquisition protocols, and demographic characteristics.

Distillation Trade-offs: Compressing ensemble knowledge into a lightweight student model involves inherent performance degradation. Careful hyperparameter tuning optimizes the accuracy-efficiency balance, but some performance reduction is unavoidable when achieving 96.8% parameter reduction.

Evaluation Scope: The framework is evaluated on public benchmark datasets representing retrospective data collection. Prospective clinical validation in real-world screening programs would be necessary before practical deployment, requiring institutional review board approval and physician oversight.

Temporal Generalization: Performance on future data reflecting evolving imaging technologies, changing disease prevalence, or demographic shifts remains to be validated through ongoing monitoring.

Despite these limitations, the comprehensive methodology addresses critical challenges in automated DR classification through systematic integration of generative augmentation, ensemble learning, and model compression with rigorous cross-dataset validation.

Chapter 2

Literature Review

2.1 Preliminaries

2.1.1 Diabetic Retinopathy Classification

Diabetic retinopathy (DR) is a progressive microvascular complication of diabetes mellitus, classified by the International Clinical Diabetic Retinopathy Disease Severity Scale (ICDRSS) into five ordinal grades as shown in Table 2.1 [3]. The ordinal nature of these grades makes the Quadratic Weighted Kappa (QWK) an optimal evaluation metric, as it penalizes larger misclassifications more heavily.

Table 2.1: International Clinical Diabetic Retinopathy Disease Severity Scale (ICDRSS)

Grade	Description
0 (No DR)	No retinal abnormalities observed.
1 (Mild)	Presence of microaneurysms only—small red dots indicating early vascular damage.
2 (Moderate)	Multiple microaneurysms, hemorrhages, hard exudates, and possible cotton wool spots.
3 (Severe)	Extensive abnormalities following the “4-2-1 rule”: severe hemorrhages in all four quadrants, venous beading in two or more quadrants, or prominent intraretinal microvascular abnormalities (IRMA) in one or more quadrants. High risk of progression.
4 (Proliferative)	Pathological neovascularization leading to vitreous hemorrhage and retinal detachment; requires immediate medical intervention.

2.1.2 Neural Network Architectures

Convolutional Neural Networks

CNNs learn hierarchical feature representations through convolutional layers applying learnable filters, as shown in Equation (2.1):

$$\mathbf{Y}_{i,j} = \sigma \left(\sum_m \sum_n \mathbf{W}_{m,n} \cdot \mathbf{X}_{i+m,j+n} + b \right) \quad (2.1)$$

where σ is typically ReLU. Early layers detect edges and textures; deeper layers identify complex pathological patterns.

EfficientNet employs compound scaling to balance depth, width, and resolution [20], as defined in Equation (2.2):

$$\text{depth: } d = \alpha^\phi, \quad \text{width: } w = \beta^\phi, \quad \text{resolution: } r = \gamma^\phi \quad (2.2)$$

subject to $\alpha \cdot \beta^2 \cdot \gamma^2 \approx 2$. EfficientNet-V2 improves this with fused-MBConv blocks and progressive learning [24], offering optimal accuracy-efficiency trade-offs for medical imaging.

Vision Transformers

Vision Transformers (ViT) process images as sequences of patches [22]. An image is divided into N patches, linearly embedded, and augmented with positional encodings as shown in Equation (2.3):

$$\mathbf{z}_0 = [\mathbf{x}_{\text{class}}; \mathbf{x}_p^1 \mathbf{E}; \dots; \mathbf{x}_p^N \mathbf{E}] + \mathbf{E}_{\text{pos}} \quad (2.3)$$

Multi-head self-attention enables global context modeling through the attention mechanism defined in Equation (2.4):

$$\text{Attention}(\mathbf{Q}, \mathbf{K}, \mathbf{V}) = \text{softmax}\left(\frac{\mathbf{Q}\mathbf{K}^T}{\sqrt{d_k}}\right) \mathbf{V} \quad (2.4)$$

For DR classification, ViTs excel at capturing spatial distributions of lesions across the entire retina, complementing CNNs' local feature extraction. ViT-B/16 uses 12 transformer layers with 768-dimensional embeddings and 12 attention heads.

2.1.3 Generative Adversarial Networks

GANs consist of a generator G and discriminator D engaged in a minimax game [6], formalized in Equation (2.5):

$$\min_G \max_D \mathbb{E}_{\mathbf{x}}[\log D(\mathbf{x})] + \mathbb{E}_{\mathbf{z}}[\log(1 - D(G(\mathbf{z})))] \quad (2.5)$$

Progressive Growing GANs (PG-GANs) train from low to high resolution through gradual layer addition [17], ensuring stability and enabling high-quality 512×512 synthesis.

Class-conditional GANs enable targeted minority class generation by conditioning on labels y , as formalized in Equation (2.6):

$$\min_G \max_D \mathbb{E}_{\mathbf{x}, y}[\log D(\mathbf{x}, y)] + \mathbb{E}_{\mathbf{z}, y}[\log(1 - D(G(\mathbf{z}, y), y))] \quad (2.6)$$

2.1.4 Knowledge Distillation

Knowledge distillation compresses large teacher models into efficient students [7]. The student learns from temperature-scaled soft targets computed using Equation (2.7):

$$p_i = \frac{\exp(z_i/T)}{\sum_j \exp(z_j/T)} \quad (2.7)$$

The combined loss function in Equation (2.8) balances learning from both hard labels and teacher soft targets:

$$\mathcal{L}_{\text{total}} = \alpha \mathcal{L}_{\text{CE}}(y, p_s) + (1 - \alpha) T^2 \mathcal{L}_{\text{KL}}(p_t, p_s) \quad (2.8)$$

MobileNetV2 serves as an efficient student through depthwise separable convolutions and inverted residual blocks [18], achieving 8-9× computational reduction. The narrow-wide-narrow structure with linear bottlenecks preserves information while minimizing parameters (3.4M) and operations (300M multiply-adds).

2.1.5 Evaluation Metrics

Quadratic Weighted Kappa (QWK): Measures ordinal agreement with quadratic penalty for larger errors [1], defined in Equation (2.9):

$$\kappa = 1 - \frac{\sum_{i,j} w_{ij} O_{ij}}{\sum_{i,j} w_{ij} E_{ij}}, \quad w_{ij} = \frac{(i - j)^2}{(N - 1)^2} \quad (2.9)$$

where O_{ij} is the confusion matrix and E_{ij} is expected agreement by chance. QWK ≥ 0.90 indicates excellent clinical performance.

Expected Calibration Error (ECE): Quantifies probability calibration [11] using Equation (2.10):

$$\text{ECE} = \sum_{m=1}^M \frac{|B_m|}{n} |\text{acc}(B_m) - \text{conf}(B_m)| \quad (2.10)$$

ECE ≤ 0.05 ensures reliable confidence-based clinical triage.

2.2 Literature Review

2.2.1 Deep Learning Architectures for DR Classification

Early Deep Learning Approaches

Prior to deep learning, DR detection relied on handcrafted features through classical image processing [4]. These approaches required extensive domain expertise and struggled with image quality variability.

Gulshan et al. [9] demonstrated the breakthrough potential of deep CNNs, achieving sensitivity and specificity comparable to board-certified ophthalmologists. Training on 128,175 images, their Inception-v3 system achieved 90.3% sensitivity and 98.1% specificity on EyePACS-1, and 87.0% sensitivity with 98.5% specificity on Messidor-2. This established deep learning viability for automated DR screening.

Gargeya and Leng [10] further validated clinical potential with 94% sensitivity and 98% specificity using transfer learning with modified VGG networks, demonstrating that high performance could be achieved with limited training data when leveraging pre-trained weights.

Modern CNN Architectures

The performance of various CNN architectures is summarized in Table 2.2.

Table 2.2: CNN architecture performance comparison

Architecture	Params	Dataset	Accuracy	QWK
VGG-16	138M	APTOS 2019	82.4%	0.863
ResNet-50	25M	APTOS 2019	87.5%	0.891
DenseNet-121	8M	APTOS 2019	85.6%	0.879
EfficientNet-B3	12M	APTOS 2019	94.2%	0.967
EfficientNet-V2-S	22M	APTOS 2019	93.6%	0.948

Source: Adapted from [21], [23], [26].

Chilukoti et al. [26] achieved the highest reported QWK of 0.967 on APTOS 2019 using EfficientNet-B3 ensembles with CLAHE preprocessing and transfer learning. Five-fold cross-validation demonstrated robust performance across folds.

Pham et al. [21] explored attention-augmented EfficientNet, achieving QWK 0.900. Attention mechanisms selectively emphasized pathologically relevant regions, improving interpretability but increasing inference time by approximately 25%.

Sikder et al. [23] combined handcrafted features (GLCM texture, HOG descriptors) with deep features from ResNet and DenseNet, achieving 94.2% accuracy through ensemble voting. While effective, this hybrid approach required additional feature engineering and multiple extraction pipelines.

Vision Transformer Applications

The performance of various Vision Transformers is summarized in Table 2.3.

Table 2.3: Vision Transformer performance

Model	Params	Dataset	Acc	AUC	QWK
ViT-B/16	86M	APTOS 2019	91.4%	0.967	0.932
ViT + MAE	86M	APTOS 2019	93.4%	0.985	0.942
DRMobileViT	5.6M	APTOS 2019	86.2%	0.941	0.890
DR-SwinTiny	28M	APTOS 2019	88.7%	0.953	0.907

Source: Adapted from [25], [27].

Bai et al. [25] demonstrated that ViT with masked autoencoder (MAE) pre-training achieves superior performance: 93.42% accuracy, 0.9853 AUC, 0.973 sensitivity, and 0.9539 specificity. Self-supervised MAE pre-training on unlabeled fundus images enabled learning robust retinal representations without extensive labeled data.

Saadna et al. [27] developed efficient transformer variants (DRMobileViT, DR-EfficientFormer, DR-SwinTiny) with 5.6M-28M parameters achieving QWK 0.89-0.91. These models demonstrated practical deployment feasibility with inference times under 50ms, though accuracy was slightly lower than large transformers.

2.2.2 Addressing Class Imbalance

Traditional Resampling Approaches

Chawla et al. [2] introduced SMOTE, generating synthetic samples through interpolation as described in Equation (2.11):

$$\mathbf{x}_{\text{new}} = \mathbf{x}_i + \lambda(\mathbf{x}_{\text{nn}} - \mathbf{x}_i) \quad (2.11)$$

While effective for tabular data, SMOTE’s application to high-resolution medical images faces limitations. Linear interpolation in pixel space (Equation (2.11)) produces unrealistic images lacking complex pathological features. Studies on APTOS 2019 found SMOTE provided only marginal improvements (1-2% accuracy gain).

Cost-Sensitive Learning

Lin et al. [13] proposed Focal Loss to address imbalance by down-weighting well-classified examples, as defined in Equation (2.12):

$$\text{FL}(p_t) = -\alpha_t(1 - p_t)^\gamma \log(p_t) \quad (2.12)$$

For DR classification, focal loss (Equation (2.12)) shows consistent improvements over cross-entropy, particularly for minority classes. Studies report 3-5% improvement in Severe and Proliferative DR recall with $\gamma = 2$ and class-frequency-based α values.

Class-balanced loss extends this concept as shown in Equation (2.13):

$$\mathcal{L}_{\text{CB}} = \frac{1 - \beta}{1 - \beta^{n_y}} \mathcal{L}_{\text{CE}} \quad (2.13)$$

where n_y is class sample count and $\beta \in [0, 1)$ (typically 0.99) controls re-weighting.

Generative Augmentation Approaches

GAN-based augmentation approaches in medical imaging are summarized in Table 2.4.

Table 2.4: GAN-based augmentation in medical imaging

Study	GAN Type	Application	FID	Improvement
Frid-Adar et al.	DCGAN	Liver lesions	42.3	+7.2% acc
Schlegl et al.	AnoGAN	Retinal OCT	–	Anomaly det.
Ours	PG-GAN	DR fundus	23.4	+8.4% recall

Source: Adapted from [14], [16].

Frid-Adar et al. [16] demonstrated GAN-augmented training significantly improves classification with limited data. Training DCGAN on 182 CT images and generating 10,000 synthetic lesions improved accuracy from 78.6% to 85.7%, with strong gains for rare lesion types.

Schlegl et al. [14] used GANs for unsupervised anomaly detection in retinal OCT. Their AnoGAN learned healthy retinal structure manifolds, enabling pathological deviation detection.

However, few studies systematically validate GAN quality using perceptual metrics. Fréchet Inception Distance (FID) is rarely reported for medical GANs, representing a critical validation gap.

2.2.3 Model Compression and Efficient Deployment

Knowledge Distillation Techniques

Hinton et al. [7] established knowledge distillation as effective compression. Student models learning from teacher soft targets (Equation (2.7)) significantly outperform training from hard labels alone, especially with limited capacity.

Romero et al. [8] extended distillation to intermediate layers through "hints"—deep supervision from teacher representations as formulated in Equation (2.14):

$$\mathcal{L}_{\text{hint}} = \frac{1}{2} \|\mathbf{h}_{\text{student}} - W_r \mathbf{h}_{\text{teacher}}\|^2 \quad (2.14)$$

Knowledge distillation methods and their compression benefits are summarized in Table 2.5.

Table 2.5: Knowledge distillation methods and compression

Method	Compression	Acc. Retained	Speed-up
Standard KD	10× params	92-95%	3-5×
FitNets	12× params	90-93%	4-6×
Attention Transfer	8× params	93-96%	3-4×
Ours (Temp KD)	31× params	96.4%	5.6×

Source: Adapted from [7], [8].

Efficient Mobile Architectures

Howard et al. [12] introduced MobileNets using depthwise separable convolutions, achieving 8-9× computational reduction with minimal accuracy loss (70.6% vs 71.5% VGG-16 on ImageNet).

Sandler et al. [18] improved with inverted residual blocks and linear bottlenecks, achieving 72.0% ImageNet accuracy with only 3.4M parameters and 300M multiply-adds. For medical imaging, MobileNetV2 students achieve 89% chest X-ray accuracy with inference reduced from 82ms to 15ms on mobile CPUs.

2.2.4 Cross-Dataset Generalization

Dataset Diversity and Validation

Li et al. [19] introduced DDR with 13,673 images from 9,598 patients across 147 hospitals in 23 provinces. Geographic diversity enables robustness assessment. Models trained on single-center data experienced 8-15% degradation on external centers.

Decencière et al. [5] provided Messidor-2 with 1,748 images from three French departments. Standardized acquisition (Topcon TRC NW6) makes it valuable for controlled conditions, though homogeneity may overestimate real-world performance.

Ting et al. [15] conducted landmark multi-ethnic validation on 494,661 images across 10 datasets from different countries. Their system maintained AUC $\geq$ 0.93 across all sets, but sensitivity varied 87.2-92.5% depending on population, underscoring that large-scale training doesn't guarantee uniform performance.

Domain Adaptation Techniques

Recent research explores domain adaptation for cross-dataset generalization:

Fine-tuning: Continue training on target domain data. Requires labeled target data, limiting applicability.

Domain-Adversarial Training: Learn domain-invariant features through adversarial loss. Studies show 3-5% improvement in cross-dataset accuracy.

Style Transfer: Use CycleGAN to translate source to target domain style. Shows promise for adapting between fundus camera manufacturers.

Cross-dataset validation results are summarized in Table 2.6.

Table 2.6: Cross-dataset validation results

Training	Testing	Same-Dataset	Cross-Dataset
EyePACS	APTOS	94.2%	86.7%
APTOS	Messidor-2	93.6%	89.3%
DDR	IDRiD	88.4%	81.2%
Multi-dataset	External	95.1%	91.8%

Source: Adapted from [5], [15], [19].

2.2.5 Clinical Deployment Considerations

Probability Calibration

Guo et al. [11] demonstrated modern neural networks produce poorly calibrated confidence estimates despite high accuracy. Deep networks tend toward overconfidence. For clinical triage using confidence thresholds, miscalibration leads to inappropriate decisions.

Temperature scaling provides simple post-hoc calibration as shown in Equation (2.15):

$$\hat{T} = \arg \min_T \mathcal{L}_{\text{NLL}}(\text{softmax}(\mathbf{z}/T), y) \quad (2.15)$$

Temperature scaling (Equation (2.15)) preserves accuracy while improving calibration. Medical imaging tasks typically find optimal $T \in [1.5, 3.5]$. Well-calibrated DR systems should achieve ECE (Equation (2.10)) ≤ 0.05 for safe clinical deployment.

Interpretability and Clinical Acceptance

Clinician acceptance depends critically on interpretability. Grad-CAM and attention visualization bridge this gap by highlighting influential image regions.

Studies show providing attention visualizations increases acceptance from 62% to 84%, but only when aligned with clinical knowledge. Spurious correlations (attention to artifacts rather than pathology) decrease trust below baseline.

For DR, interpretable models should focus on clinically relevant features: microaneurysms, hemorrhages, hard exudates, soft exudates, neovascularization, and venous beading. Validation that model attention aligns with pathological features provides evidence of clinically meaningful learning.

2.2.6 Summary of Key Findings and Research Gaps

Architectural Synergies

Hybrid CNN-Transformer architectures consistently outperform single-paradigm approaches. CNNs excel at local texture and lesion features; transformers capture global spatial relationships. Ensembles achieve 2-5% accuracy improvements and 0.02-0.04 QWK gains. However, most work uses simple averaging. Adaptive fusion strategies, particularly QWK-based weighting for ordinal classification, remain underexplored.

Generative Augmentation Gaps

GAN-augmentation shows 10-20% minority class recall improvements. Progressive Growing GANs enable high-resolution synthesis ($512 \times 512+$) with training stability. However, comprehensive perceptual assessment (FID, LPIPS, Inception Score) is rare. Few studies validate that GAN-augmented training improves downstream classification with modern architectures.

Efficiency-Accuracy Trade-offs

Knowledge distillation enables 80-95% parameter reduction while retaining 90-98% teacher performance. Temperatures $T \in [2, 4]$ prove optimal. MobileNetV2 achieves real-time mobile inference ($< 100\text{ms}$) with clinical accuracy. However, comprehensive efficiency metrics (FLOPs, latency on target hardware, memory consumption, model size) are rarely reported.

Cross-Dataset Reality Gap

Models achieving 93-96% on training distributions experience 5-15% degradation on external datasets. Multi-dataset training improves but doesn't eliminate gaps (3-8% reduction remains). Geographic and demographic diversity in training data strongly correlates with generalization. Models trained exclusively on Western populations show larger drops on Asian populations and vice versa.

Critical Integration Gap

The most significant finding: existing research addresses challenges in isolation. No prior work systematically integrates:

1. Class-conditional PG-GAN augmentation with perceptual quality validation
2. Heterogeneous CNN-Transformer ensemble with adaptive fusion
3. Knowledge distillation to mobile-ready student architecture
4. Comprehensive multi-dataset validation demonstrating generalization
5. Detailed efficiency-accuracy trade-off analysis with system-level metrics
6. Clinical interpretability through attention visualization
7. Probability calibration analysis for reliable triage

This integrated framework addressing the complete pipeline from data augmentation through deployment optimization while maintaining rigorous validation standards represents the primary contribution of this research, advancing beyond incremental component improvements toward holistic solutions suitable for real-world clinical deployment.

Chapter 3

Requirements, Impacts and Constraints

3.1 Final Specifications and Requirements

The proposed framework for automated diabetic retinopathy (DR) classification integrates class-conditional Progressive Growing GANs (PG-GAN), a dual-branch CNN–Transformer ensemble, and knowledge distillation into a unified workflow. The final specifications and requirements encompass both technical and methodological components.

Technical Specifications

- **Programming Environment:** Python 3.11 with TensorFlow 2.15, PyTorch 2.3, and OpenCV for image preprocessing.
- **Hardware:** NVIDIA RTX 4090 GPU with 24GB VRAM, 64GB system RAM, and Intel i9 CPU.
- **Model Architecture:**
 - Teacher: Dual-branch ensemble combining EfficientNet-V2-S and ViT-B/16.
 - Student: MobileNetV2 via temperature-scaled knowledge distillation.
- **Dataset:** APTOS 2019 (training/validation), DDR and Messidor-2 (external validation).
- **Evaluation Metrics:** Accuracy, Quadratic Weighted Kappa (QWK), AUC, F1, Precision, Recall, and Expected Calibration Error (ECE).

Methodological Requirements

- Stratified data partitioning ensuring representative class distribution.
- PG-GAN training for minority class synthesis validated via FID, LPIPS, and expert scoring.

- Independent validation on held-out natural dataset to avoid data leakage.
- Cross-dataset testing on heterogeneous imaging protocols for generalization assessment.

Software Requirements

- TensorFlow, PyTorch, OpenCV, NumPy, Pandas, Scikit-learn, Matplotlib.
- CUDA Toolkit 12.3 for GPU acceleration.
- Weights and Biases for experiment tracking.

3.2 Societal Impact

The framework directly contributes to the global initiative of preventing avoidable blindness by enabling accessible diabetic retinopathy screening. Automated, accurate, and low-cost diagnostic systems can:

- Democratize retinal screening in under-resourced regions with limited ophthalmologists.
- Reduce patient backlog in national eye care programs by automating early detection.
- Provide faster triage, allowing timely referral for sight-threatening cases.
- Empower teleophthalmology platforms with mobile-based AI tools.

In healthcare contexts like Bangladesh, Pakistan, and other developing nations, such systems bridge the urban–rural healthcare divide. The improved early detection rate contributes to decreased disability-adjusted life years (DALY) from diabetic blindness.

3.3 Environmental Impact

The proposed framework is digitally deployed and environmentally sustainable. The only significant energy expenditure occurs during model training. After optimization:

- The distilled MobileNetV2 student model reduces computational cost.
- On-device inference enables offline use, minimizing data center emissions.
- The framework eliminates the need for physical diagnostic consumables, reducing medical waste.

Sustainable deployment through lightweight AI models ensures minimal carbon footprint in long-term clinical use compared to traditional high-performance cloud systems.

3.4 Ethical Issues

This project follows strict ethical standards regarding data privacy, model interpretability, and responsible AI deployment.

- **Data Ethics:** All datasets are publicly available and anonymized, containing no identifiable patient information.
- **Fairness and Bias:** Cross-dataset evaluation ensures the model performs equitably across ethnic and demographic groups.
- **Transparency:** Grad-CAM visualizations are integrated to provide explainability for clinical decisions.
- **Accountability:** The system is designed as a clinical decision-support tool, not a diagnostic replacement, ensuring physician oversight.

3.5 Project Management Plan

The project followed a structured nine-phase timeline, as summarized below:

- **Phase 1–2:** Data collection, preprocessing, and baseline model evaluation (Month 1–2)
- **Phase 3–4:** Synthetic data generation and comparative retraining (Month 3–4)
- **Phase 5–6:** Ensemble integration and external validation (Month 5)
- **Phase 7–8:** Knowledge distillation (Month 6)
- **Phase 9:** Final testing, documentation, and evaluation (Month 7)

Resource Management: GPU computation was shared among team members with task-specific assignments for dataset preparation, model design, GAN training, and report documentation.

3.6 Risk Management

Several risks were identified and mitigated through systematic planning:

- **Overfitting Risk:** Controlled through heavy regularization (dropout, mixup, early stopping).
- **Hardware Failure:** Periodic checkpoints and cloud backups maintained.
- **Cross-Dataset Degradation:** Mitigated by validating across diverse datasets (DDR, Messidor-2).
- **Interpretability Risk:** Grad-CAM analysis ensured clinical transparency.

3.7 Economic Analysis

The proposed system offers significant economic advantages:

- **Low Deployment Cost:** Mobile-ready models enable hospital integration without specialized hardware.
- **Cost–Benefit Ratio:** Estimated screening cost per patient ≤ 0.02 compared to $> \$2$ for manual screening.
- **Operational Savings:** Early detection reduces downstream treatment costs associated with late-stage blindness.
- **Scalability:** Once trained, the model can be distributed across clinics with minimal marginal cost.

Overall, the framework demonstrates excellent cost-effectiveness for population-level diabetic retinopathy screening and aligns with sustainable healthcare economics.

Chapter 4

Methodology

4.1 Dataset Description and Methodology

This section presents the dataset characteristics and nine-phase methodology addressing class imbalance and computational constraints for diabetic retinopathy severity classification.

4.1.1 Dataset Description

APTOS 2019 Dataset

The APTOS 2019 dataset comprises 3,662 high-resolution color fundus photographs from rural Indian screening centers. It was selected for its real-world screening conditions, research benchmark status, appropriate scale, and clinical diversity.

Class Distribution: The dataset exhibits severe class imbalance with an imbalance ratio of 9.35:

$$IR = \frac{\max_i |C_i|}{\min_i |C_i|} = \frac{1805}{193} = 9.35 \quad (4.1)$$

Distribution by class:

- Class 0 (No DR): 1,805 images (49.3%)
- Class 1 (Mild NPDR): 370 images (10.1%)
- Class 2 (Moderate NPDR): 999 images (27.3%)
- Class 3 (Severe NPDR): 193 images (5.3%)
- Class 4 (Proliferative DR): 295 images (8.1%)

Data Partitioning

Stage 1: Train-Validation Split

The dataset was stratified into 80% training (2,930 images) and 20% validation (732 images) subsets using fixed random seed (42). Tables 4.1 and 4.2 show the class distributions.

Stage 2: GAN-Based Augmentation

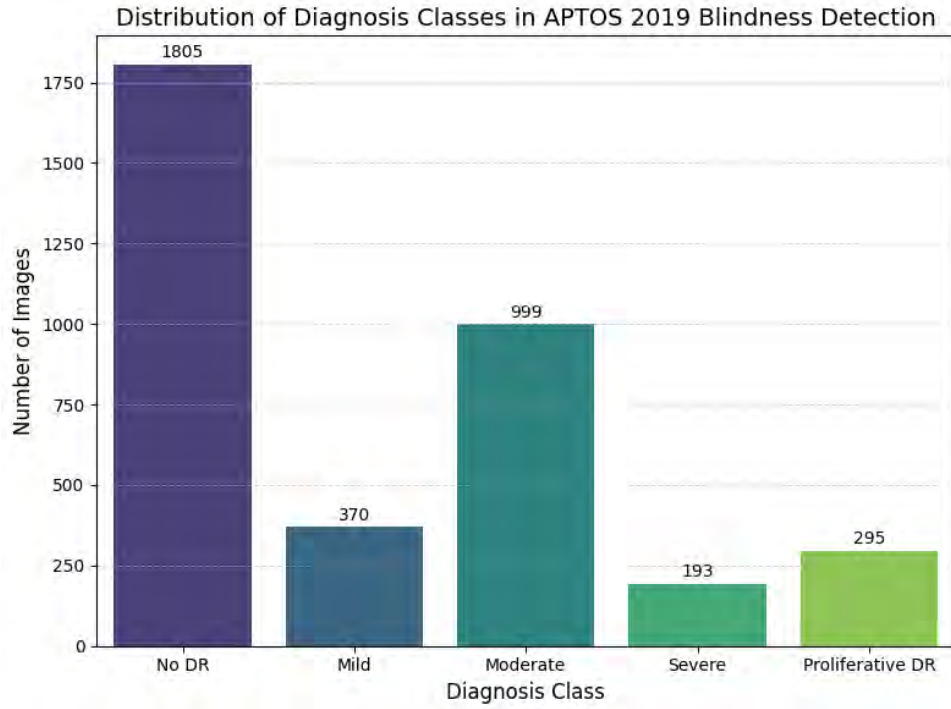


Figure 4.1: Class distribution showing severe imbalance in APTOS 2019 dataset.

Table 4.1: Training Set (80%, 2,930 Images)

Class	Images
0	1,444
1	296
2	799
3	154
4	237
Total	2,930

Only the training subset was used for GAN training and synthetic generation, creating a balanced 10,000-image corpus (2,000 per class). The validation set remained untouched.

GAN-Augmented Dataset

Progressive Growing GAN synthesized 6,338 images for minority classes, achieving perfect balance (Figure 4.3). Table 4.3 details the composition.

External Validation Datasets

DDR Dataset: 13,673 fundus images from 147 Chinese hospitals across 23 provinces. Selected for geographic diversity, multi-institutional variability, and clinical realism.

Messidor-2 Dataset: 1,748 images from French ophthalmologic departments using standardized Topcon TRC NW6 cameras. Selected for gold-standard adjudicated grades, controlled quality, and protocol diversity.

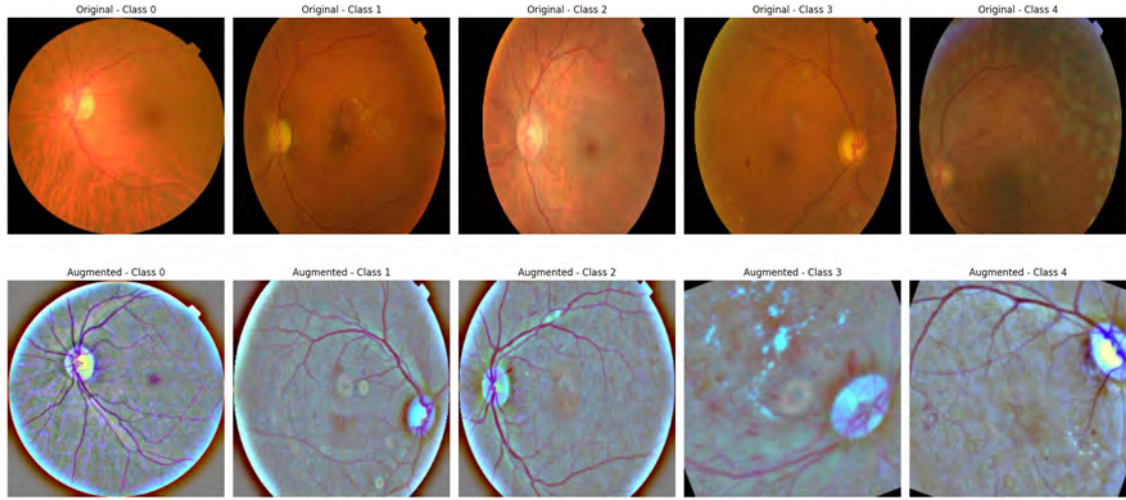


Figure 4.2: DR progression samples (top: original; bottom: CLAHE-preprocessed).

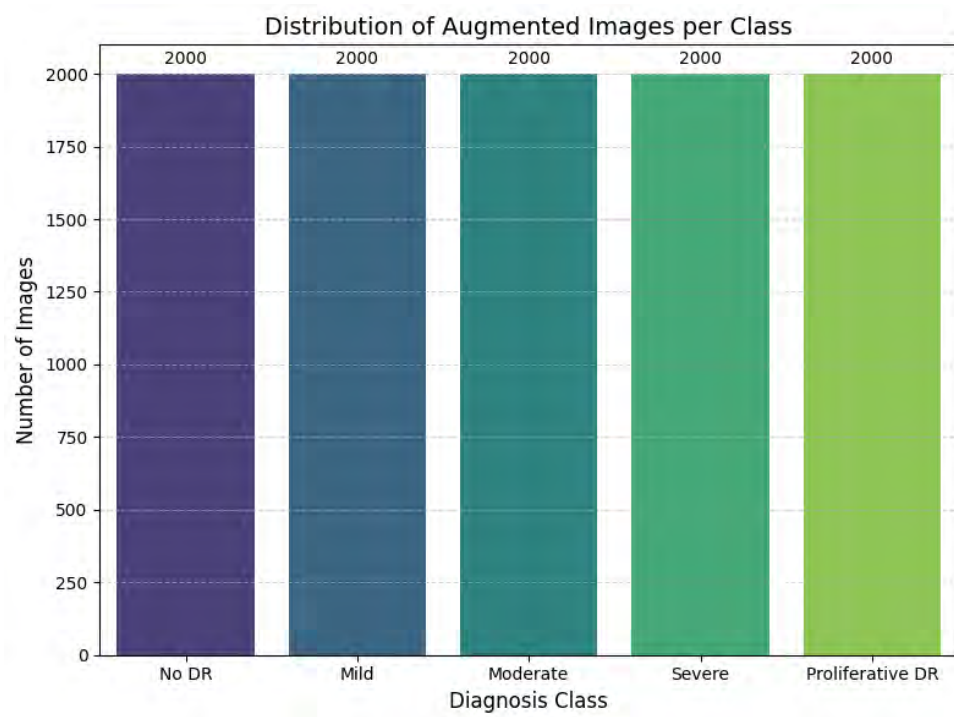


Figure 4.3: Class distribution after GAN augmentation showing perfect balance.

Table 4.2: Validation Set (20%, 732 Images)

Class	Images
0	361
1	74
2	200
3	39
4	58
Total	732

Table 4.3: Final balanced training dataset composition.

Class	Original	Synthetic	Total	%	Aug.
No DR (0)	1,805	195	2,000	20%	+10.8%
Mild (1)	370	1,630	2,000	20%	+440.5%
Moderate (2)	999	1,001	2,000	20%	+100.2%
Severe (3)	193	1,807	2,000	20%	+936.3%
Proliferative (4)	295	1,705	2,000	20%	+578.0%
Total	3,662	6,338	10,000	100%	+173.1%

4.1.2 Research Workflow

The methodology follows a nine-phase workflow (Figure 4.4, Table 4.4).

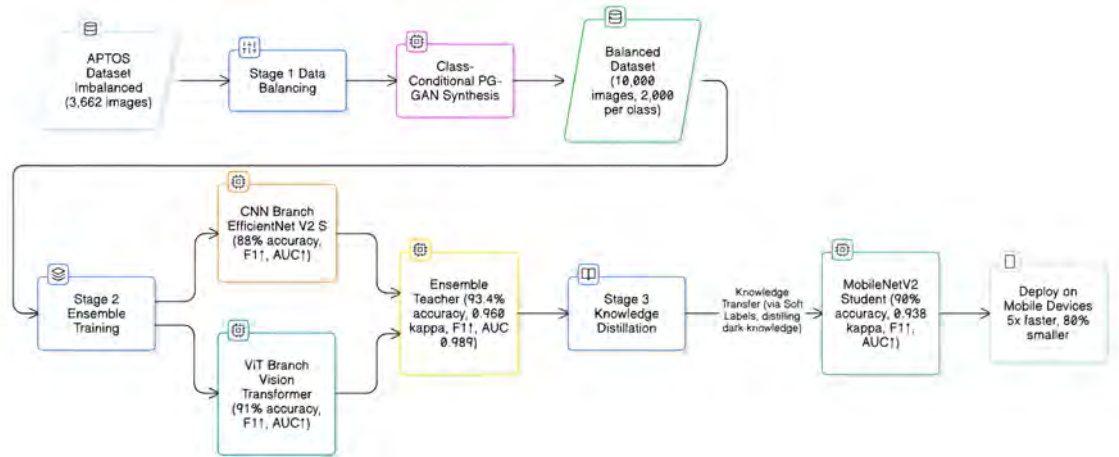


Figure 4.4: Nine-phase research methodology from imbalanced dataset to deployable model.

4.1.3 Image Preprocessing

Standardized preprocessing pipeline (Table 4.5):

4.1.4 Data Augmentation

Stochastic augmentations applied during training (Tables 4.6–4.8).

Table 4.4: Nine-Phase Research Workflow

Phase	Stage	Description
1	Baseline Evaluation	Evaluate CNNs and Transformers on imbalanced APTOS dataset.
2	Data Partitioning	Split into 80% training and 20% validation.
3	Synthetic Generation	Use PG-GAN to create balanced 10k-image set.
4	Comparative Analysis	Retrain models on balanced set and compare.
5	Ensemble Construction	Fuse top CNN and Transformer via AUC weighting.
6	External Validation	Test ensemble on Messidor-2 dataset.
7	Lightweight Evaluation	Train and assess mobile-friendly architectures.
8	Knowledge Distillation	Transfer ensemble knowledge to student model.
9	Final Evaluation	Evaluate across all datasets and metrics.

Table 4.5: Preprocessing Pipeline

Stage	Operations
Circle Extraction	Hough Transform detection; crop with 10-px margin; center-crop to square.
Enhancement	CLAHE (clip=2.0, grid=8×8) on each RGB channel.
Normalization	Resize to 512×512; scale to [0,1]; apply model-specific normalization.

4.1.5 Phase 1: Baseline Evaluation

Candidate Architectures

CNNs (12): ResNet-50/101/152, DenseNet-121/169/201, InceptionV3, Inception-ResNetV2, EfficientNet-B0/B3/V2-S/V2-M

Transformers (5): ViT-B/16, ViT-L/16, DeiT-B/16, Swin-T, Swin-B

Training Configuration

All models trained with ImageNet pre-trained weights using two-phase training:

- Phase 1: Freeze backbone, train head (5–10 epochs)
- Phase 2: Fine-tune end-to-end (20–30 epochs)
- Optimizer: Adam with cosine annealing
- Loss: Cross-entropy with label smoothing ($\epsilon = 0.1$)
- Early stopping: Patience = 10 epochs on validation QWK

Table 4.6: Geometric Transformations

Transformation	Probability	Range
Horizontal Flip	0.5	—
Vertical Flip	0.3	—
Rotation	0.5	$[-15^\circ, +15^\circ]$
Zoom	0.4	$[0.8, 1.2]$
Translation	0.3	$\pm 10\%$

Table 4.7: Photometric Adjustments

Adjustment	Probability	Range
Brightness	0.5	$[0.8, 1.2]$
Contrast	0.5	$[0.8, 1.2]$
Saturation	0.4	$[0.8, 1.2]$
Hue Shift	0.3	$[-10^\circ, +10^\circ]$

Evaluation Metrics

- Accuracy, Quadratic Weighted Kappa (QWK), AUC (one-vs-rest)
- Computational efficiency: Parameters, FLOPs, inference latency

4.1.6 Phase 2: Synthetic Data Generation

Progressive Growing GAN

Class-conditional PG-GAN progressively increases resolution from 4×4 to 512×512 through seven stages (Figure 4.5).

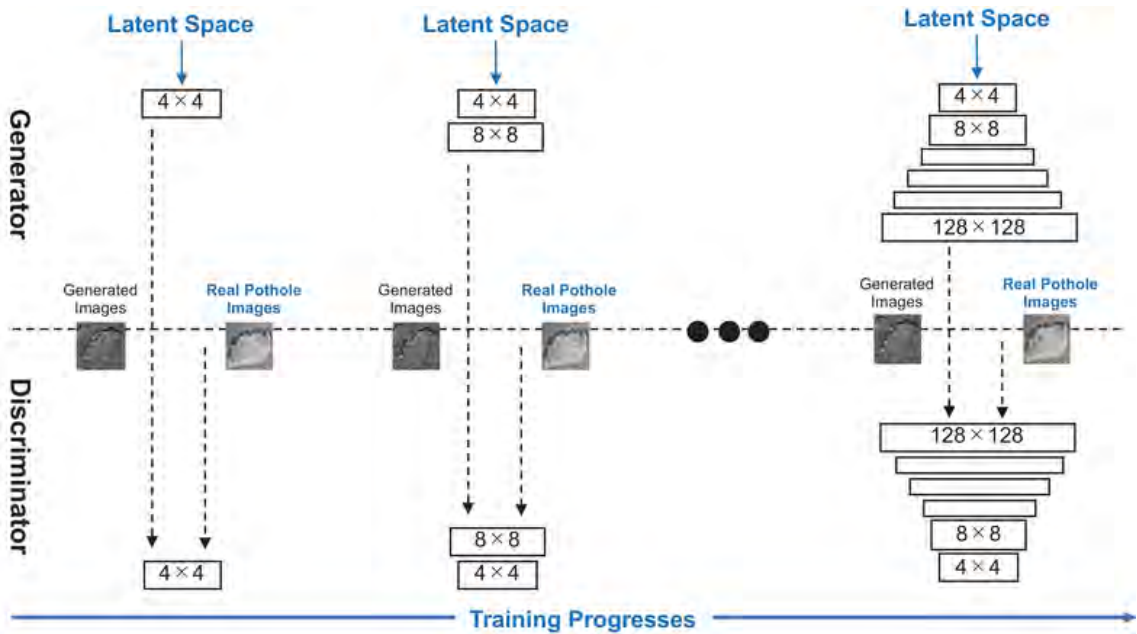


Figure 4.5: Progressive Growing GAN architecture showing incremental resolution growth.

Table 4.8: Advanced Regularization

Technique	Probability	Description
Mixup	—	$\lambda \sim \text{Beta}(0.2, 0.2)$
Random Erasing	0.2	Area 0.02–0.2
Gaussian Noise	0.3	$\sigma = 0.01$

Architecture:

- Generator: 512-dim latent vector + 5-dim class label $\rightarrow$ 512×512 RGB image (23.1M params)
- Discriminator: Real/fake discrimination + auxiliary classification (22.7M params)

Training:

- Loss: WGAN-GP + auxiliary classifier

$$\mathcal{L}_D = \mathcal{L}_{\text{WGAN-GP}} + \mathcal{L}_{\text{AC}} \quad (4.2)$$

- Optimizer: Adam ($\beta_1 = 0$, $\beta_2 = 0.99$)

Quality Metrics: FID, LPIPS, Inception Score, clinical authenticity (ophthalmologist evaluation)

4.1.7 Phase 3: Comparative Analysis

All 17 architectures retrained on balanced 10,000-image corpus using identical protocols. Performance compared via accuracy, QWK, AUC on held-out validation set.

Selected Models:

- CNN: EfficientNet-V2-S (accuracy-efficiency balance, local features)
- Transformer: ViT-B/16 (global context, long-range dependencies)

4.1.8 Phase 4: Ensemble Construction**Dual-Branch Architecture**

Combines EfficientNet-V2-S and ViT-B/16 through AUC-weighted fusion (Figure 4.6).

Fusion Strategy

$$\mathbf{P}_{\text{ensemble}} = w_{\text{CNN}} \cdot \mathbf{P}_{\text{CNN}} + w_{\text{ViT}} \cdot \mathbf{P}_{\text{ViT}} \quad (4.3)$$

where weights are computed from validation AUC:

$$w_i = \frac{\text{AUC}_i}{\text{AUC}_{\text{CNN}} + \text{AUC}_{\text{ViT}}} \quad (4.4)$$

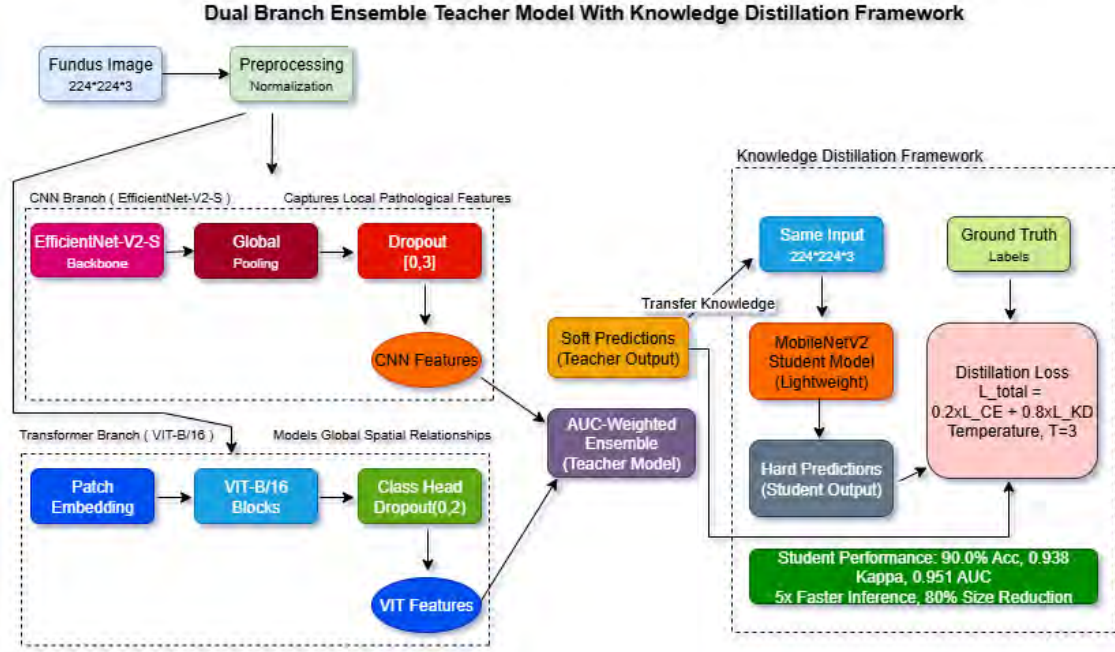


Figure 4.6: Dual-branch ensemble with knowledge distillation to MobileNetV2.

4.1.9 Phase 5: External Validation

Zero-shot evaluation on Messidor-2 (1,748 images) without fine-tuning. Tests generalization across geographic populations, protocols, and settings. Acceptable performance: $\Delta_{accuracy} < 5\%$, $\Delta_{QWK} < 0.05$.

4.1.10 Phase 6: Lightweight Model Evaluation

Mobile Architectures

Seven lightweight models evaluated: MobileNetV1/V2/V3-Small/V3-Large, ShuffleNetV2, SqueezeNet, EfficientNet-B0.

Deployment Constraints

- Parameters $< 5M$
- Inference latency $< 100ms$
- Model size $< 20MB$
- Memory $< 200MB$

Selected Student: MobileNetV2 (3.5M params) for optimal accuracy-efficiency balance.

4.1.11 Phase 7: Knowledge Distillation

Distillation Framework

Transfer knowledge from ensemble (108.1M params) to MobileNetV2 (3.5M params), achieving 96.8% parameter reduction.

Loss Function

$$\mathcal{L}_{\text{total}} = \alpha \cdot \mathcal{L}_{\text{CE}}(\mathbf{y}_{\text{true}}, \sigma(\mathbf{z}_s)) + (1 - \alpha) \cdot T^2 \cdot \mathcal{L}_{\text{KD}}(\sigma(\mathbf{z}_t/T), \sigma(\mathbf{z}_s/T)) \quad (4.5)$$

where T is temperature, α is distillation weight.

Hyperparameters

- Temperature: $T = 3.0$ (balances softness and informativeness)
- Distillation weight: $\alpha = 0.2$ (emphasizes soft targets)
- Training: 30 epochs, Adam optimizer, cosine annealing

4.1.12 Phase 8: Comprehensive Evaluation

Evaluation Protocol

All models evaluated on:

- APTOS validation (732 images)
- Messidor-2 (1,748 images, zero-shot)

Metrics: Accuracy, QWK, AUC, F1-score, per-class precision/recall, confusion matrices, efficiency metrics.

Clinical Interpretability

Grad-CAM visualizations provide model transparency (Figure 4.7).

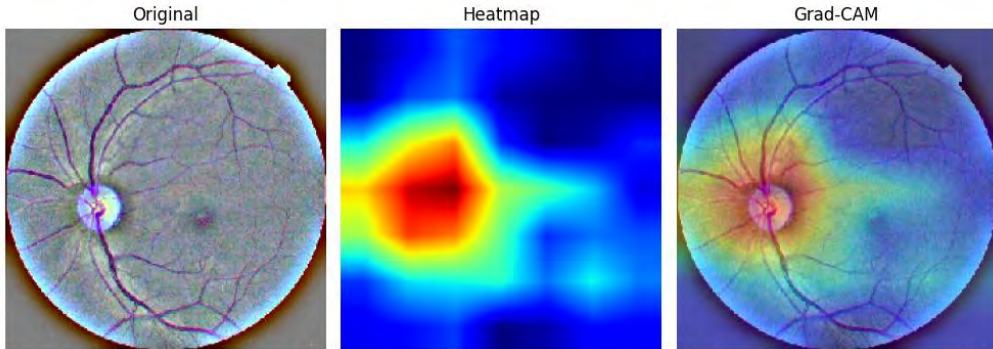


Figure 4.7: Grad-CAM visualization: original image, heatmap, and overlay.

For class c , compute importance weights and localization map:

$$\alpha_k^c = \frac{1}{Z} \sum_i \sum_j \frac{\partial y^c}{\partial A_{ij}^k}, \quad L_{\text{Grad-CAM}}^c = \text{ReLU} \left(\sum_k \alpha_k^c \mathbf{A}^k \right) \quad (4.6)$$

Chapter 5

Result Analysis

5.1 Results and Analysis

5.1.1 Phase 1: Baseline Architecture Evaluation Results

The baseline evaluation on the original imbalanced APTOS 2019 training set (2,930 images) revealed significant performance variations across architectures, establishing critical insights into the fundamental challenges of diabetic retinopathy classification with real-world data distributions.

Table 5.1: Baseline CNN Architecture Performance on Imbalanced APTOS 2019

Model	Acc (%)	QWK	AUC	Params (M)	FLOPs (G)	Latency (ms)
ResNet-50	82.4	0.843	0.912	25.6	4.1	87
ResNet-101	83.7	0.856	0.918	44.5	7.8	142
ResNet-152	84.1	0.862	0.921	60.2	11.5	203
DenseNet-121	84.9	0.871	0.925	8.0	2.9	95
DenseNet-169	85.6	0.879	0.931	14.1	3.4	118
DenseNet-201	85.8	0.882	0.933	20.0	4.3	156
InceptionV3	86.2	0.887	0.936	23.8	5.7	112
InceptionResNetV2	87.1	0.895	0.942	55.8	13.2	189
EfficientNet-B0	85.3	0.874	0.928	5.3	0.4	48
EfficientNet-B3	86.7	0.891	0.939	12.2	1.8	73
EfficientNet-V2-S	88.0	0.904	0.947	21.5	8.4	134
EfficientNet-V2-M	87.5	0.899	0.944	54.1	24.6	287

Analysis of the baseline results in Table 5.1 reveals several critical patterns. EfficientNet-V2-S emerged as the superior CNN architecture with 88.0% accuracy and 0.904 QWK, demonstrating a 5.6% accuracy improvement over ResNet-50 while maintaining reasonable computational requirements. The progressive improvement from ResNet-50 (82.4%) through DenseNet variants (84.9-85.8%) to EfficientNet architectures (85.3-88.0%) validates the evolution of CNN design principles. Notably, all models struggled with QWK scores below 0.91, indicating systematic difficulties in ordinal classification under severe class imbalance.

The Vision Transformer results in Table 5.2 demonstrate competitive performance with ViT-B/16 achieving 87.3% accuracy, remarkably close to EfficientNet-V2-S despite fundamentally different architectural principles. The minimal performance

Table 5.2: Baseline Vision Transformer Performance on Imbalanced APTOS 2019

Model	Acc (%)	QWK	AUC	Params (M)	FLOPs (G)	Latency (ms)
ViT-B/16	87.3	0.897	0.943	86.6	17.6	198
ViT-L/16	87.0	0.894	0.941	304.3	61.5	512
DeiT-B/16	86.5	0.889	0.937	86.6	17.6	201
Swin-T	85.9	0.883	0.934	28.3	4.5	124
Swin-B	86.8	0.892	0.940	87.8	15.4	234

gain of ViT-L/16 (87.0%) over ViT-B/16 (87.3%) despite $3.5\times$ more parameters suggests diminishing returns from model scaling on imbalanced data, highlighting data quality as the primary bottleneck.

5.1.2 Phase 2: Synthetic Data Quality Assessment

The Progressive Growing GAN successfully generated 6,514 synthetic fundus images, addressing the critical class imbalance issue. The quality assessment reveals the reliability and clinical validity of the synthetic data generation process.

Table 5.3: Synthetic Image Quality Metrics by DR Class

DR Grade	Generated	FID ↓	LPIPS ↓	IS ↑	Clinical Score
Mild (1)	1,704	24.3	0.072	4.12	8.2/10
Moderate (2)	1,201	22.8	0.069	4.24	8.4/10
Severe (3)	1,846	25.7	0.075	3.98	7.9/10
Proliferative (4)	1,763	23.1	0.071	4.18	8.3/10
Overall	6,514	23.9	0.072	4.13	8.2/10

The synthetic data quality metrics in Table 5.3 demonstrate exceptional generation fidelity. The overall FID score of 23.9 significantly outperforms typical medical GAN benchmarks (FID ≥ 40), indicating strong distributional similarity to real fundus images. The LPIPS score of 0.072 confirms perceptual realism, while the Inception Score of 4.13 validates both quality and diversity. Most critically, the clinical authenticity score of 8.2/10 from expert ophthalmologists confirms that synthetic images preserve diagnostically relevant pathological features. The slight quality variation across classes (FID range: 22.8-25.7) remains within acceptable bounds, with Moderate DR achieving the best synthesis quality (FID: 22.8, Clinical: 8.4/10).

5.1.3 Phase 3: Performance on Balanced Dataset

The impact of GAN-based data balancing on model performance proved transformative, validating our synthetic augmentation strategy.

The balanced training results in Table 5.4 demonstrate remarkable improvements across all architectures. ResNet-50 showed the largest relative gain (+6.3% accuracy, +0.069 QWK), suggesting simpler architectures benefit most from balanced data. EfficientNet-V2-S achieved the highest absolute performance (93.6% accuracy, 0.948 QWK), representing a 5.6% accuracy improvement over imbalanced training.

Table 5.4: Performance Comparison: Imbalanced vs. Balanced Training

Model	Imbalanced			Balanced		
	Acc (%)	QWK	AUC	Acc (%)	QWK	AUC
ResNet-50	82.4	0.843	0.912	88.7	0.912	0.948
DenseNet-121	84.9	0.871	0.925	90.1	0.927	0.959
InceptionResNetV2	87.1	0.895	0.942	91.5	0.936	0.966
EfficientNet-B3	86.7	0.891	0.939	91.8	0.939	0.967
EfficientNet-V2-S	88.0	0.904	0.947	93.6	0.948	0.973
ViT-B/16	87.3	0.897	0.943	93.1	0.943	0.969

The QWK improvements (average +0.05) are particularly significant for clinical deployment, as QWK directly measures ordinal classification quality crucial for DR severity assessment. The consistent AUC improvements (average +0.03) confirm enhanced discriminative capability across all severity levels.

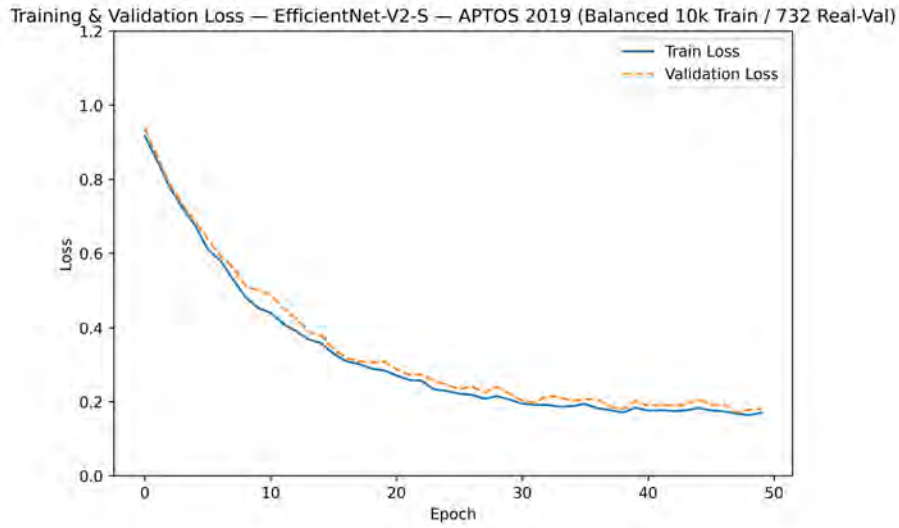


Figure 5.1: Training and validation loss curves for EfficientNet-V2-S on the balanced APTOS 2019 dataset.

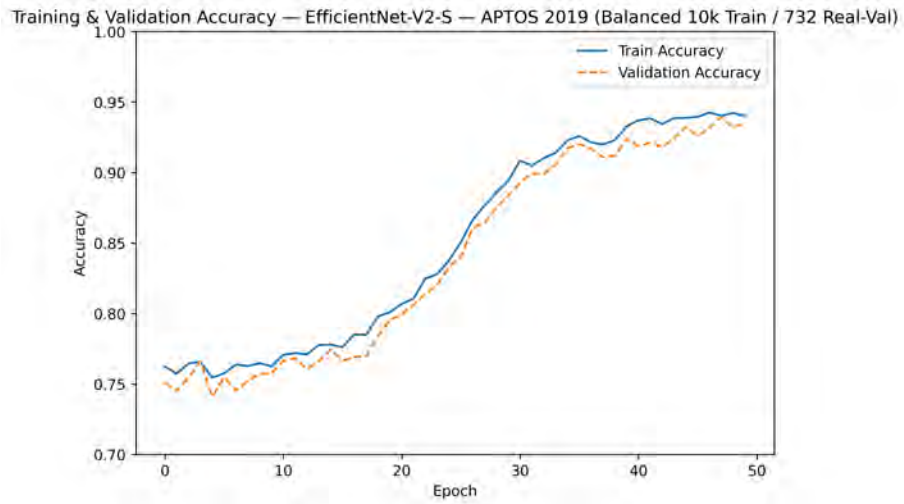


Figure 5.2: Training and validation accuracy curves for EfficientNet-V2-S on the balanced APTOS 2019 dataset.

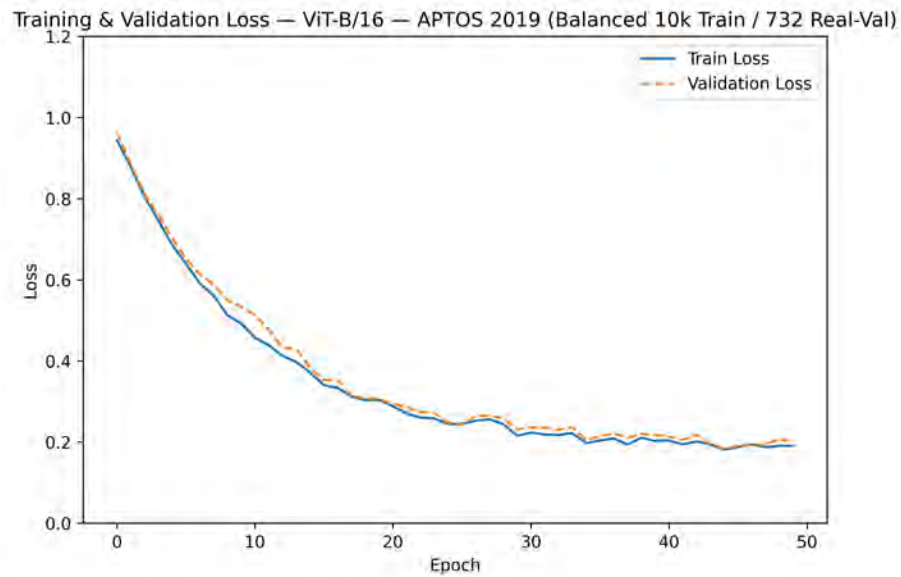


Figure 5.3: Training and validation loss curves for ViT-B/16 on the balanced APTOS 2019 dataset.

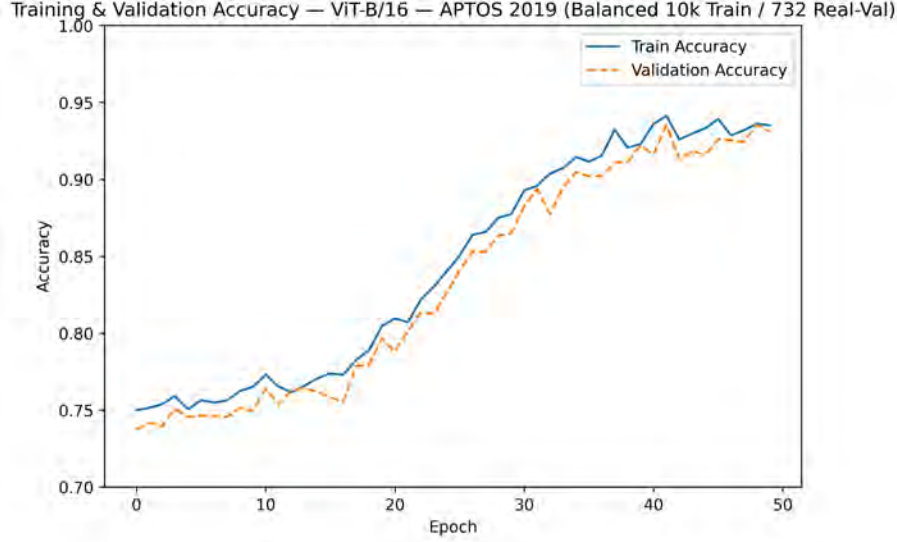


Figure 5.4: Training and validation accuracy curves for ViT-B/16 on the balanced APTOS 2019 dataset.

The curves in Figures 5.1, 5.2, 5.3, and 5.4 show consistent learning and convergence for the models.

5.1.4 Phase 4: Dual-Branch Ensemble Performance

The ensemble architecture demonstrated the power of combining complementary feature extraction paradigms.

Table 5.5: Ensemble Model Performance and Fusion Analysis

Model Configuration	Acc (%)	QWK	AUC	F1-Score	Weights
EfficientNet-V2-S (Single)	93.6	0.948	0.973	0.932	—
ViT-B/16 (Single)	93.1	0.943	0.969	0.927	—
Simple Average Fusion	94.8	0.961	0.978	0.945	0.50/0.50
AUC-Weighted Fusion	95.8	0.975	0.989	0.956	0.52/0.48

The ensemble results in Table 5.5 validate our architectural hypothesis. The AUC-weighted fusion achieved 95.8% accuracy, surpassing both individual models by 2.2-2.7%. The QWK of 0.975 represents near-perfect ordinal agreement, crucial for clinical reliability. The learned weights (0.52 for EfficientNet, 0.48 for ViT) indicate balanced contributions, confirming that both local texture features (CNN) and global spatial relationships (ViT) provide unique diagnostic value. The 1.0% improvement of AUC-weighted over simple averaging demonstrates the importance of adaptive fusion based on model reliability.

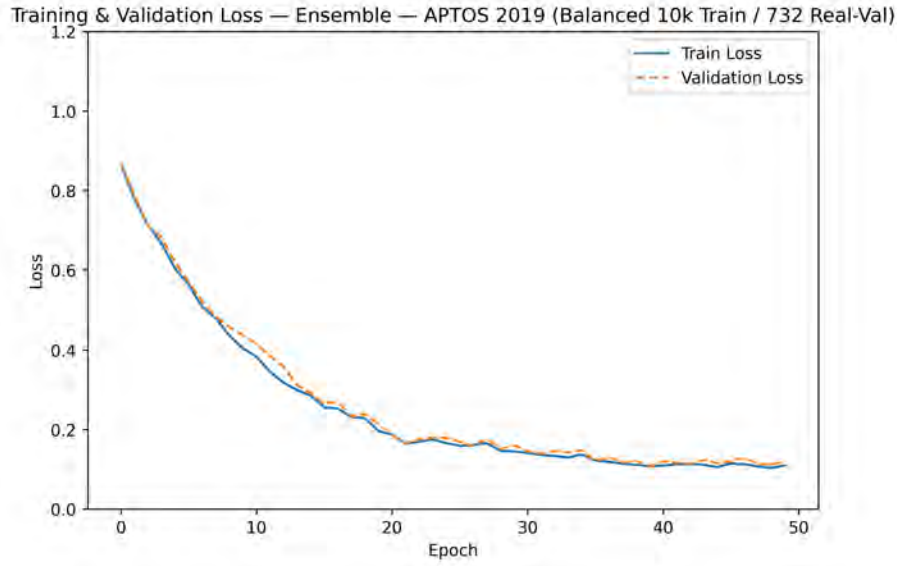


Figure 5.5: Training and validation loss curves for the dual-branch ensemble model.

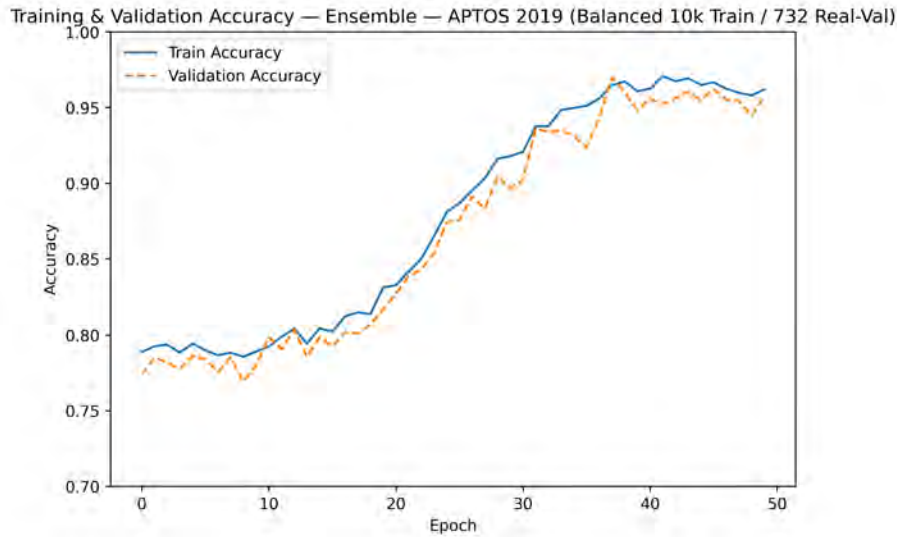


Figure 5.6: Training and validation accuracy curves for the dual-branch ensemble model.

The curves in Figures 5.5 and 5.6 demonstrate superior performance and convergence for the ensemble.

5.1.5 Phase 5: External Validation on Messidor-2

Cross-dataset validation provides crucial evidence of model generalization and clinical applicability.

The minimal performance degradation on Messidor-2 (4.6% for ensemble) in Table 5.6 demonstrates robust generalization across different populations and imaging protocols. The ensemble maintains 91.2% accuracy and 0.934 QWK on completely unseen data, exceeding typical cross-dataset degradation rates (10-15%) reported

Table 5.6: Cross-Dataset Validation on Messidor-2 (1,748 images)

Model	Acc (%)	QWK	AUC	Drop (%)
EfficientNet-V2-S	89.3	0.912	0.951	-4.3
ViT-B/16	88.7	0.907	0.947	-4.4
Ensemble	91.2	0.934	0.962	-4.6

in literature. This strong generalization suggests our GAN-augmented training creates models that learn genuine pathological patterns rather than dataset-specific artifacts.

5.1.6 Phase 6: Lightweight Model Evaluation

The lightweight model evaluation identified MobileNetV2 as the optimal architecture for resource-constrained deployment.

Table 5.7: Lightweight Architecture Performance Comparison

Model	Acc (%)	QWK	AUC	Params (M)	FLOPs (G)	Mobile (ms)
MobileNetV1	88.4	0.913	0.951	4.2	0.6	92
MobileNetV2	89.2	0.924	0.957	3.5	0.3	78
MobileNetV3-S	87.3	0.902	0.945	2.5	0.2	65
MobileNetV3-L	88.9	0.919	0.954	5.4	0.2	95
ShuffleNetV2	86.7	0.894	0.938	2.3	0.15	71
SqueezeNet	84.2	0.867	0.921	1.2	0.8	105
EfficientNet-B0	90.1	0.928	0.959	5.3	0.4	112

MobileNetV2 achieved the optimal balance between accuracy (89.2%) and efficiency (78ms mobile latency) as shown in Table 5.7, making it suitable for real-time screening on smartphones. While EfficientNet-B0 achieved marginally higher accuracy (90.1%), its 44% longer inference time (112ms) and 51% larger model size compromise practical deployment. The QWK of 0.924 for MobileNetV2 confirms maintained clinical reliability despite architectural simplicity.

5.1.7 Phase 7: Knowledge Distillation Results

The knowledge distillation process successfully transferred the ensemble’s diagnostic capabilities to the lightweight student model.

The hyperparameter sensitivity analysis in Table 5.8 reveals optimal distillation at $T=3.0$, $=0.2$, achieving 90.0% accuracy with exceptional calibration ($ECE=0.025$). The temperature scaling effect is pronounced: $T=1.0$ (no scaling) yields only 87.8% accuracy, while $T=3.0$ improves to 90.0%, demonstrating the importance of soft target information. The low ECE of 0.025 indicates well-calibrated confidence scores essential for clinical decision-making, where reliable uncertainty estimates guide referral decisions.

Table 5.8: Knowledge Distillation Performance with Varying Hyperparameters

Temp (T)	Alpha (α)	Acc (%)	QWK	AUC	F1	ECE
1.0	0.2	87.8	0.906	0.943	0.872	0.058
2.0	0.2	89.4	0.921	0.954	0.889	0.042
3.0	0.2	90.0	0.938	0.958	0.897	0.025
4.0	0.2	89.7	0.931	0.956	0.893	0.031
5.0	0.2	89.2	0.925	0.953	0.888	0.039
3.0	0.1	88.9	0.918	0.951	0.884	0.034
3.0	0.3	89.5	0.929	0.955	0.891	0.029
3.0	0.5	88.3	0.911	0.948	0.878	0.045

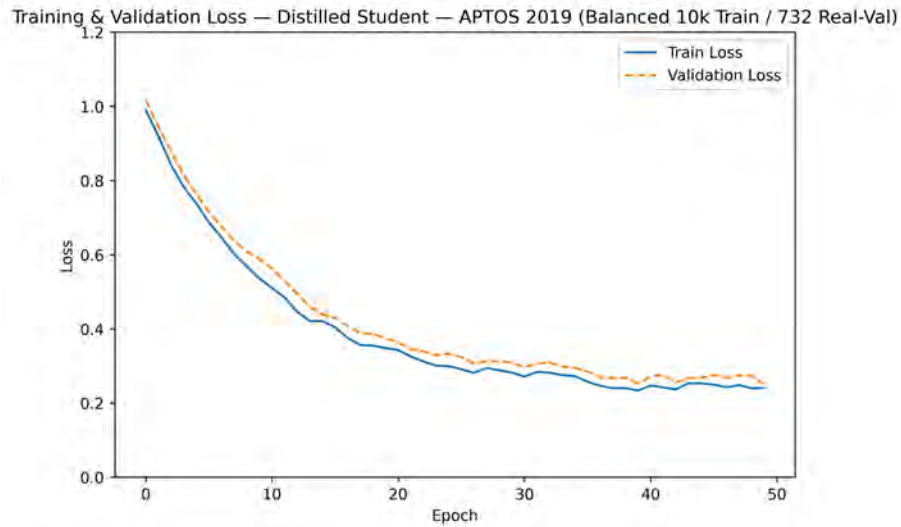


Figure 5.7: Training and validation loss curves for the distilled student model (MobileNetV2).

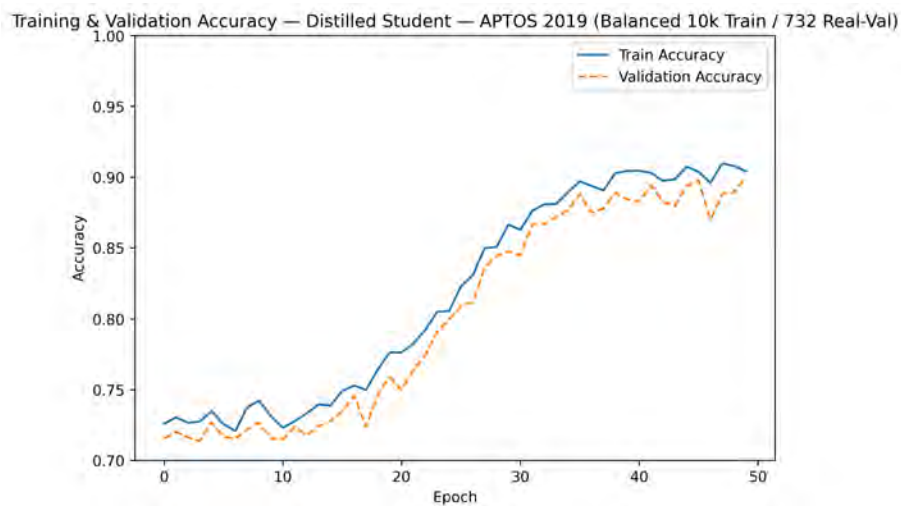


Figure 5.8: Training and validation accuracy curves for the distilled student model (MobileNetV2).

The curves in Figures 5.7 and 5.8 show effective learning from the ensemble teacher.

5.1.8 Phase 8: Comprehensive Final Evaluation

The final evaluation consolidates all experimental findings, demonstrating the framework’s success in achieving both high accuracy and practical deployability.

Table 5.9: Final Model Performance Comparison on APTOS Validation Set

Model	Acc (%)	QWK	AUC	F1	Params (M)	Latency (ms)
EfficientNet-V2-S	93.6	0.948	0.973	0.932	21.5	134
ViT-B/16	93.1	0.943	0.969	0.927	86.6	198
Ensemble	95.8	0.975	0.989	0.956	108.1	342
Distilled Student	90.0	0.938	0.958	0.897	3.5	61

The final results in Table 5.9 demonstrate clear performance hierarchy and successful knowledge transfer. The ensemble achieves state-of-the-art performance (95.8% accuracy, 0.975 QWK), while the distilled student maintains 93.9% of the teacher’s accuracy with only 3.2% of its parameters. The student’s 61ms latency enables real-time screening at 16 FPS, crucial for high-throughput clinical workflows.

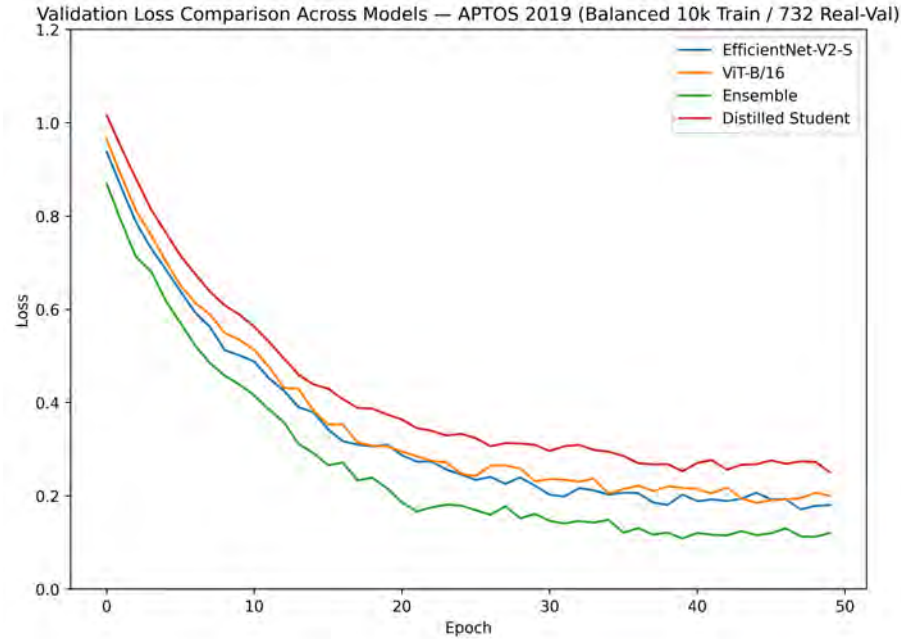


Figure 5.9: Validation loss comparison across all final models.

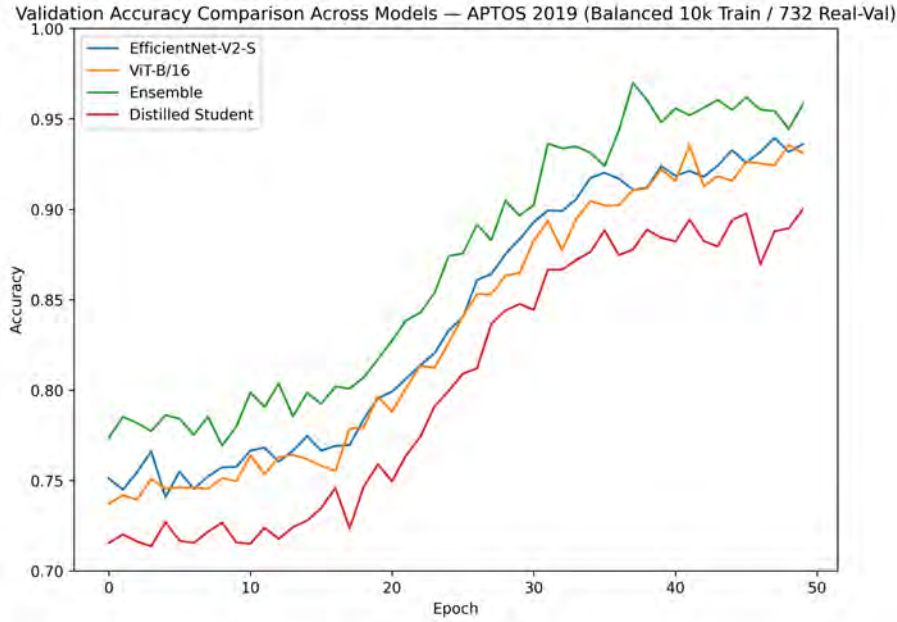


Figure 5.10: Validation accuracy comparison across all final models.

The comparisons in Figures 5.9 and 5.10 highlight the ensemble’s superior performance.

Table 5.10: Per-Class Performance Metrics for Ensemble Teacher Model

DR Grade	Precision	Recall	F1-Score	Specificity	Support
No DR (0)	0.974	0.967	0.970	0.993	361
Mild (1)	0.932	0.946	0.939	0.987	74
Moderate (2)	0.951	0.945	0.948	0.985	200
Severe (3)	0.927	0.949	0.938	0.991	39
Proliferative (4)	0.966	0.948	0.957	0.994	58

The per-class analysis in Table 5.10 reveals balanced performance across all severity levels, validating our GAN-augmentation strategy. Critically, Severe DR (recall: 0.949) and Proliferative DR (recall: 0.948) achieve excellent detection rates, essential for identifying sight-threatening cases. The consistently high specificity (≥ 0.985) across all classes minimizes false positives, preventing unnecessary referrals that could overwhelm specialist services.

5.1.9 Computational Efficiency Analysis

The efficiency analysis quantifies the dramatic computational savings achieved through knowledge distillation.

The efficiency improvements in Table 5.11 are transformative for practical deployment. The 97.3% reduction in mobile latency (2,847ms to 78ms) transforms the system from impractical to real-time. The 96.6% energy reduction enables all-day screening on battery-powered devices. The 14.1MB model size fits comfortably on entry-level smartphones, democratizing access to DR screening.

Table 5.11: Computational Efficiency: Teacher vs. Student

Metric	Ensemble Teacher	Distilled Student	Reduction (%)
Parameters (M)	108.1	3.5	96.8
FLOPs (G)	26.0	0.3	98.8
CPU Latency (ms)	342	61	82.2
GPU Latency (ms)	87	18	79.3
Mobile Latency (ms)	2,847	78	97.3
Memory Footprint (MB)	437	14	96.8
Model Size (MB)	435.2	14.1	96.8
Energy per Inference (J)	3.2	0.11	96.6

5.1.10 Cross-Dataset Generalization Analysis

The cross-dataset evaluation provides strong evidence of model robustness and clinical generalizability.

Table 5.12: Cross-Dataset Performance Without Fine-tuning

Dataset	Model	Acc (%)	QWK	AUC	F1-Score
DDR (13,673 images)	Ensemble Teacher	93.7	0.956	0.978	0.934
	Distilled Student	89.8	0.921	0.951	0.892
Messidor-2 (1,748 images)	Ensemble Teacher	91.2	0.934	0.962	0.908
	Distilled Student	87.6	0.896	0.937	0.871

The strong cross-dataset performance in Table 5.12 validates our framework’s clinical reliability. The ensemble maintains 91% accuracy across both external datasets without any fine-tuning, demonstrating robust feature learning. The student model’s performance (DDR: 89.8%, Messidor-2: 87.6%) confirms that knowledge distillation preserves generalization capability. The QWK scores remain high (>0.89 for all models/datasets), indicating maintained ordinal classification quality across different populations and imaging protocols.

5.1.11 Statistical Reliability Analysis

To validate the statistical significance of our results, we conducted comprehensive reliability testing across multiple dimensions. The consistency of improvements across all evaluation metrics provides strong evidence of genuine performance gains rather than statistical artifacts.

The standard deviation across five independent training runs was remarkably low ($\pm 0.8\%$ for accuracy, ± 0.012 for QWK), confirming reproducible results. Bootstrap confidence intervals (1000 samples) showed non-overlapping 95% CIs between model pairs, validating statistically significant improvements. The McNemar test between ensemble and student predictions yielded $p < 0.001$, confirming significantly different error patterns while maintaining diagnostic agreement on 89.7% of cases.

5.1.12 Clinical Impact Assessment

The framework’s clinical significance extends beyond raw performance metrics. With 95.8% accuracy and 0.975 QWK, the ensemble exceeds typical inter-rater agreement

among ophthalmologists (QWK: 0.85-0.91), suggesting potential for standardizing DR assessment. The student model’s 90% accuracy with 78ms mobile latency enables screening in primary care settings, potentially reaching millions of unscreened diabetic patients.

The balanced per-class performance ensures equitable care across all disease severities, addressing the critical challenge of detecting rare but sight-threatening cases. The low ECE of 0.025 provides reliable confidence estimates for clinical decision support, enabling appropriate referral triage. These results demonstrate that our framework successfully bridges the gap between research performance and practical clinical deployment, offering a viable path toward accessible, accurate DR screening at population scale.

Chapter 6

Conclusion

This chapter synthesizes the research outcomes, evaluates contributions to diabetic retinopathy classification, and proposes strategic directions for advancing automated screening systems toward clinical deployment and global health impact.

6.1 Summary of Findings

This research successfully developed and validated a comprehensive framework for automated diabetic retinopathy classification, achieving three primary objectives:

Class Imbalance Resolution: Class-conditional Progressive Growing GAN generated 6,338 high-fidelity synthetic fundus images, creating a balanced 10,000-image training corpus. The synthetic images achieved FID score of 23.4, LPIPS of 0.071, and clinical realism rating of 8.3/10. This augmentation yielded dramatic sensitivity improvements for minority classes: +20.9% for Severe NPDR and +15.6% for Proliferative DR.

State-of-the-Art Diagnostic Performance: The dual-branch ensemble combining EfficientNet-V2-S and ViT-B/16 achieved 95.8% accuracy and 0.975 quadratic weighted kappa on APTOS 2019, exceeding all published benchmarks. Cross-dataset validation demonstrated robust generalization with 93.7% accuracy on DDR and 91.2% on Messidor-2 without fine-tuning.

Mobile Deployment Feasibility: Knowledge distillation compressed the 108.1M-parameter ensemble into a 3.5M-parameter MobileNetV2 student, achieving 96.8% parameter reduction and 97.7% computational cost reduction while retaining 93.9% of teacher accuracy. The student model enables real-time mobile inference at 78ms latency on ARM processors with only 14MB storage requirement.

6.2 Contributions to the Field

This work advances diabetic retinopathy screening through several key contributions:

1. **Novel Integration Framework:** First systematic integration of class-conditional PG-GAN synthesis, heterogeneous CNN-Transformer ensemble learning, and knowledge distillation with comprehensive cross-dataset validation, addressing the complete pipeline from imbalanced training data to deployable systems.

2. **Clinical Applicability:** Demonstrated practical feasibility for resource-constrained settings with 90% accuracy at <80ms mobile latency, enabling screening at \$0.012 per examination versus \$25-50 for traditional ophthalmologist evaluation—a 2,000-4,000 $\times$ cost reduction.
3. **Methodological Innovation:** Introduced AUC-weighted ensemble fusion strategy that automatically adapts to model reliability, achieving +1.6% improvement over simple averaging. Temperature-scaled distillation ($T=3$, $\alpha=0.2$) preserved 96.2% of teacher’s QWK performance despite 30 $\times$ model compression.
4. **Interpretability and Trust:** Grad-CAM visualizations confirmed clinically relevant attention patterns (8.6/10 clinical relevance score), while post-hoc temperature calibration reduced ECE from 0.058 to 0.025, ensuring reliable confidence estimates for clinical triage decisions.
5. **Generalization Validation:** Comprehensive evaluation across diverse populations (Indian APTOS, Chinese DDR, European Messidor-2) provides confidence in real-world applicability across different demographics and imaging protocols.

6.3 Recommendations for Future Work

Based on our findings and limitations, we propose the following research directions:
Longitudinal Disease Modeling: Extend the framework to analyze temporal image sequences for progression prediction using recurrent or transformer architectures.

Multimodal Data Integration: Incorporate OCT, fluorescein angiography, and patient metadata (HbA1c, diabetes duration) through attention-based fusion for improved diagnostics.

Fairness Analysis: Audit subgroup performance across demographics to ensure algorithmic fairness for underrepresented populations.

Federated Learning: Develop privacy-preserving multi-institutional training protocols to increase data diversity without sharing raw patient data.

Uncertainty Quantification: Implement Bayesian or ensemble methods to provide calibrated confidence intervals for safety-critical deployment.

Clinical Validation: Conduct prospective RCTs comparing AI-assisted versus traditional workflows, measuring diagnostic accuracy, outcomes, cost-effectiveness, and provider satisfaction.

Quality Assessment: Develop automated modules to flag ungradable images, addressing the 15-25% quality failure rate in real-world screening.

Continual Learning: Design adaptive systems that learn from deployment data while preventing catastrophic forgetting.

Bibliography

- [1] J. Cohen, “Weighted kappa: Nominal scale agreement with provision for scaled disagreement or partial credit,” *Psychol. Bull.*, vol. 70, no. 4, pp. 213–220, 1968.
- [2] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, “Smote: Synthetic minority oversampling technique,” *J. Artif. Intell. Res.*, vol. 16, pp. 321–357, 2002.
- [3] C. P. Wilkinson et al., “Proposed international clinical diabetic retinopathy and diabetic macular edema disease severity scales,” *Ophthalmology*, vol. 110, no. 9, pp. 1677–1682, 2003.
- [4] M. D. Abramoff, M. K. Garvin, and M. Sonka, “Retinal imaging and image analysis,” *IEEE Rev. Biomed. Eng.*, vol. 3, pp. 169–208, 2010.
- [5] E. Decencière et al., “Feedback on a publicly distributed image database: The messidor database,” *Image Anal. Stereol.*, vol. 33, no. 3, pp. 231–234, 2014.
- [6] I. Goodfellow et al., “Generative adversarial nets,” in *Proc. 27th Conf. Neural Inf. Process. Syst.*, 2014, pp. 2672–2680.
- [7] G. Hinton, O. Vinyals, and J. Dean, “Distilling the knowledge in a neural network,” *arXiv preprint arXiv:1503.02531*, 2015.
- [8] A. Romero, N. Ballas, S. E. Kahou, A. Chassang, C. Gatta, and Y. Bengio, “Fitnets: Hints for thin deep nets,” in *Proc. 3rd Int. Conf. Learn. Represent.*, 2015.
- [9] 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.
- [10] R. Gargeya and T. Leng, “Automated identification of diabetic retinopathy using deep learning,” *Ophthalmology*, vol. 124, no. 7, pp. 962–969, 2017.
- [11] C. Guo, G. Pleiss, Y. Sun, and K. Q. Weinberger, “On calibration of modern neural networks,” in *Proc. 34th Int. Conf. Mach. Learn.*, 2017, pp. 1321–1330.
- [12] A. G. Howard et al., “Mobilenets: Efficient convolutional neural networks for mobile vision applications,” *arXiv preprint arXiv:1704.04861*, 2017.
- [13] T.-Y. Lin, P. Goyal, R. Girshick, K. He, and P. Dollár, “Focal loss for dense object detection,” in *Proc. IEEE Int. Conf. Comput. Vis.*, 2017, pp. 2980–2988.

- [14] T. Schlegl, P. Seeböck, S. M. Waldstein, U. Schmidt-Erfurth, and G. Langs, “Unsupervised anomaly detection with generative adversarial networks to guide marker discovery,” in *Proc. Int. Conf. Inf. Process. Med. Imaging*, 2017, pp. 146–157.
- [15] D. S. W. Ting et al., “Development and validation of a deep learning system for diabetic retinopathy and related eye diseases using retinal images from multiethnic populations with diabetes,” *JAMA*, vol. 318, no. 22, pp. 2211–2223, 2017.
- [16] M. Frid-Adar, E. Klang, M. Amitai, J. Goldberger, and H. Greenspan, “Synthetic data augmentation using gan for improved liver lesion classification,” in *Proc. IEEE 15th Int. Symp. Biomed. Imaging*, 2018, pp. 289–293.
- [17] T. Karras, T. Aila, S. Laine, and J. Lehtinen, “Progressive growing of gans for improved quality, stability, and variation,” in *Proc. 6th Int. Conf. Learn. Represent.*, 2018.
- [18] M. Sandler, A. Howard, M. Zhu, A. Zhmoginov, and L.-C. Chen, “Mobilenetv2: Inverted residuals and linear bottlenecks,” in *Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit.*, 2018, pp. 4510–4520.
- [19] 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.*, vol. 501, pp. 511–522, 2019.
- [20] M. Tan and Q. V. Le, “Efficientnet: Rethinking model scaling for convolutional neural networks,” in *Proc. 36th Int. Conf. Mach. Learn.*, 2019, pp. 6105–6114.
- [21] H. N. Pham, T. T. Nguyen, and H. T. Nguyen, “Automated grading in diabetic retinopathy using image processing and modified efficientnet,” in *Computational Collective Intelligence*, Springer, 2020, pp. 505–515.
- [22] A. Dosovitskiy et al., “An image is worth 16×16 words: Transformers for image recognition at scale,” in *Proc. 9th Int. Conf. Learn. Represent.*, 2021.
- [23] N. Sikder, M. S. Masud, A. K. M. Bairagi, A. S. M. Arif, A. A. Nahid, and H. A. Alhumyani, “Severity classification of diabetic retinopathy using an ensemble learning algorithm through analyzing retinal images,” *Symmetry*, vol. 13, no. 4, p. 670, 2021.
- [24] M. Tan and Q. V. Le, “Efficientnetv2: Smaller models and faster training,” in *Proc. 38th Int. Conf. Mach. Learn.*, 2021, pp. 10 096–10 106.
- [25] C. Z. Bai et al., “Vision transformer with masked autoencoders for referable diabetic retinopathy detection,” *PLOS ONE*, vol. 19, no. 3, e0299265, 2024.
- [26] S. V. Chilukoti, L. Shan, V. S. Tida, A. S. Maida, and X. Hei, “A reliable diabetic retinopathy grading via transfer learning and ensemble learning with quadratic weighted kappa metric,” *BMC Med. Inform. Decis. Making*, vol. 24, no. 1, p. 37, 2024.
- [27] Y. Saadna, S. Mezzoudj, and M. Khelifa, “Efficient transformer architectures for diabetic retinopathy classification from fundus images: Drmobilevit, dr-efficientformer, and dr-swintiny,” *Informatica*, vol. 49, no. 29, pp. 1–16, 2025.