

ReCAN: A Light-Weight Residual Channel Attention Network with Alternating Skip Connection for Robust Medical Image Classification

by

Mohammed Abdul Al Arafat Tanzin
21301178

MD Abrar Uddin
21301159

Md. Jisan Mashrafi
21301058

Abrar Fahim
21301073

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
Summer 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:

Mohammed Abdul Al Arafat Tanzin

21301178

MD Abrar Uddin

21301159

MD. Jisan Mashrafi

21301058

Abrar Fahim

21301073

Approval

The thesis titled “ ReCAN: A Light-Weight Residual Channel Attention Network with Alternating Skip Connection for Robust Medical Image Classification ” submitted by

1. Mohammed Abdul Al Arafat Tanzin (21301178)
2. MD Abrar Uddin (21301159)
3. Md. Jisan Mashrafi (21301058)
4. Abrar Fahim (21301073)

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

Examining Committee:

Supervisor:
(Member)

Dr. Mohammad Golam Robiul Alam

Professor
Department of Computer Science and Engineering
Brac University

Co-Supervisor:
(Member)

Nirjhor Datta

Lecturer
Department of Computer Science and Engineering
BRAC University

Thesis Coordinator:
(Member)

Dr. Mohammad Golam Robiul Alam

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

Ethics Statement

Our research confirms the claim made, and all cited and referenced papers and sources are accurate. The work has never been submitted for a degree to any other college or academic organization. The four co-authors acknowledge and accept any violations of the thesis rule. In addition, we would like to take this chance to thank everyone who has supported us through this process. We finished our thesis without committing to any illegal techniques. Our work complies with the ethical standards of Brac University.

Abstract

This thesis presents ReCAN, a lightweight convolutional neural network tailored for efficient and accurate medical image classification on the MedMNIST benchmark. Motivated by the need for compact architectures that can deliver state-of-the-art results on resource-constrained hardware, this work integrates a ResNet-18 backbone with a novel multi-stage channel attention mechanism comprising five sequential layers with varying reduction ratios. The architecture is methodologically founded on iterative experimentation with dilated convolutions, advanced residual connections, and ablation studies on attention modules to optimize the trade-off between model complexity and classification performance. Comprehensive evaluations on multiple MedMNIST datasets, including PathMNIST and others, demonstrate that ReCAN consistently outperforms MedMamba—both its Tiny in terms of accuracy, while achieving lower parameter counts and floating point operations (FLOPs) sometime and larger variants as well in terms of accuracy. The results establish that careful design of channel attention and skip connections within a CNN backbone can surpass more computationally intensive models without sacrificing generalization. ReCAN thus contributes a new, efficient deep learning baseline for medical image analysis, supporting future research in scalable biomedical AI applications.

Keywords: ReCAN, Channel Attention, ResNet-18, Lightweight CNN, Medical Image Classification, MedMNIST, Deep Learning, Parameter Efficiency, FLOPs Reduction, PathMNIST, Biomedical Image Analysis, Efficient Neural Networks

Acknowledgment

First and foremost, we are thankful to the Almighty for giving us the chance and the path we needed to finish this project on time. Secondly, we would like to sincerely thank our thesis supervisor, Dr. Mohammad Golam Robiul Alam and Nirjhor Datta for his constant encouragement, mentorship, and guidance as we explored a challenging topic. Their relentless assistance and continuous input enabled us to conquer the challenges. Thirdly, we intend to use this opportunity to express our gratitude to every faculty member for their assistance and support during our stay at Brac University. Finally, we would like to thank our beloved parents for their unwavering support, encouragement, and prayers.

Table of Contents

Declaration	i
Approval	ii
Ethics Statement	iv
Abstract	v
Dedication	vi
Acknowledgment	vi
Table of Contents	vii
List of Figures	x
List of Tables	xiii
Nomenclature	xiii
1 Introduction	1
1.1 Background	1
1.2 Rationale of the Study or Motivation	1
1.2.1 Research Gap and Motivation	2
1.2.2 Identifying the Gap	2
1.2.3 Motivation for ReCAN	4
1.3 Objectives	5
1.3.1 Research Questions	5
1.4 Scopes and Challenges	6
1.5 Methodology in Brief	7
2 Literature Review	9
2.1 Preliminaries	9
2.2 Review of Existing Research	10
2.2.1 Foundational CNN Architectures in Medical Imaging	10
2.2.2 Attention Mechanisms in Medical Image Classification	10
2.2.3 Hybrid and State-Space Models for Medical Imaging	11
2.2.4 Lightweight and Efficient Architectures	12
2.2.5 Multi-Stage and Hierarchical Attention Designs	12
2.2.6 Transfer Learning and Domain Adaptation	13
2.2.7 Benchmark Studies on MedMNIST	13

2.2.8	Regularization Techniques in Medical Image Classification . .	14
2.2.9	Interpretability and Explainability in Medical AI	14
3	Requirements, Impacts and Constraints	16
3.1	Final Specifications and Requirements	16
3.2	Societal Impact	16
3.3	Environmental Impact	17
3.4	Ethical Issues	17
3.5	Risk Management	17
3.6	Economic Analysis	18
4	Proposed Methodology	19
4.1	Design Process or Methodology Overview	19
4.1.1	ReCAN Architecture Design Philosophy	21
4.1.2	ResNet-18 Backbone Architecture	21
4.1.3	ReCAN Channel Attention Stack	22
4.1.4	Global Pooling and Feature Aggregation	24
4.1.5	Complete ReCAN Architecture	25
4.2	Data Collection and Preprocessing	25
4.2.1	Dataset Source	25
4.2.2	Preprocessing Pipeline	25
4.2.3	Dataset Descriptions	26
4.3	Implementation of Selected Design	27
4.3.1	Software Framework and Development Environment	27
4.3.2	Dataset Loading and Preprocessing Implementation	27
4.3.3	Model Architecture Implementation	28
4.3.4	Evaluation Metrics Implementation	28
4.3.5	Computational Environment and Training Time	29
4.3.6	Training Configuration and Hyperparameters	29
4.3.7	Summary of Implementation	31
5	Result Analysis	32
5.1	Performance Evaluation	32
5.1.1	Evaluation Criteria	32
5.1.2	Testing Methodology	33
5.2	Analysis of Design Solutions	33
5.2.1	Feature Enhancement Branch and Main Branch	33
5.3	Final Design Adjustments	34
5.3.1	Channel Attention Layer Ablations	34
5.3.2	Skip Connection Strategies	35
5.3.3	MLP Classifier Head Refinement	36
5.4	Statistical Analysis	37
5.4.1	Performance Variability and Confidence Intervals	37
5.4.2	Training Dynamics Analysis	38
5.5	Comparisons and Relationships	38
5.5.1	Parameter Efficiency Comparison	38
5.5.2	Comparison with State-of-the-Art Models through different dataset	39
5.5.3	Comparison with Traditional Baselines	41

5.5.4	Architectural Insights and Design Trade-offs	42
5.5.5	Efficiency and Practical Deployment	42
5.6	Visual Results and Analysis	43
5.6.1	Performance Evaluation on the PathMNIST Dataset using ReCAN	43
5.6.2	Performance Evaluation on the DermaMNIST Dataset using ReCAN	47
5.6.3	Performance Evaluation on the BloodMNIST Dataset using ReCAN	51
5.6.4	Performance Evaluation on the BreastMNIST Dataset using ReCAN	55
5.6.5	Performance Evaluation on the PneumoniaMNIST Dataset using ReCAN	59
5.6.6	Performance Evaluation on the RetinaMNIST Dataset using ReCAN	63
5.7	Discussion	66
5.7.1	Overview of Findings	66
5.7.2	Architectural Simplicity as a Design Principle	67
5.7.3	Multi-Stage Channel Attention Hierarchy	67
5.7.4	The Critical Role of Skip Connections in Attention Stacks	68
5.8	Comparative Performance Analysis	69
5.8.1	ReCAN vs. Vision Transformers	69
5.8.2	ReCAN vs. Hybrid State-Space Models	70
5.8.3	Cross-Dataset Generalization Patterns	72
5.9	Model Efficiency and Interpretability	73
5.9.1	Computational Efficiency and Deployment Feasibility	73
5.9.2	Aggressive Regularization and Generalization	74
5.9.3	Attention Mechanisms and Model Interpretability	75
5.9.4	Transfer Learning and Domain Adaptation	76
5.10	Limitations and Future Scope	77
5.10.1	Benchmark Limitations and Real-World Complexity	77
5.10.2	Dataset Scale and Rare Disease Challenges	78
5.10.3	Extension to 3D Medical Imaging	78
6	Conclusion	80
6.1	Summary of Key Findings and Implications	80
6.2	Limitations	81
6.3	Future Work	82
	Bibliography	84

List of Figures

4.1	Overview of ReCAN architecture: ResNet-18 backbone followed by multi-stage channel attention blocks and MLP classifier head.	20
4.2	Backbone architecture (Resnet-18)	21
4.3	Detailed Diagram of a ReCAN Attention Block. This figure illustrates the internal structure of a single attention block used within the Channel Attention Stack.	22
5.1	Confusion matrix for PathMNIST using ReCAN showing classification performance across nine tissue pathology classes	43
5.2	Training and validation accuracy and loss curves for PathMNIST demonstrating convergence behavior	44
5.3	Correctly classified samples from PathMNIST using ReCAN showcasing diverse tissue morphologies	44
5.4	t-SNE visualization of PathMNIST features learned by ReCAN showing discriminative embedding space	45
5.5	Class distribution for PathMNIST training set revealing dataset composition and potential imbalances	46
5.6	Confusion matrix for DermaMNIST using ReCAN across seven dermatological lesion categories	47
5.7	Training and validation accuracy and loss curves for DermaMNIST showing stable optimization	48
5.8	Correctly classified samples from DermaMNIST using ReCAN demonstrating invariance to imaging variations	48
5.9	t-SNE visualization of DermaMNIST features showing semantic clustering of lesion types	49
5.10	Class distribution for DermaMNIST training set highlighting lesion category frequencies	50
5.11	Confusion matrix for BloodMNIST using ReCAN across eight blood cell types	51
5.12	Training and validation accuracy and loss curves for BloodMNIST demonstrating stable learning	52
5.13	Correctly classified samples from BloodMNIST using ReCAN showing robust cell identification	52
5.14	t-SNE visualization of BloodMNIST features revealing discriminative cell type embeddings	53
5.15	Class distribution for BloodMNIST training set showing blood cell type frequencies	54
5.16	Confusion matrix for BreastMNIST using ReCAN for binary malignancy classification	55

5.17	Training and validation accuracy and loss curves for BreastMNIST showing reliable convergence	56
5.18	Correctly classified samples from BreastMNIST using ReCAN demonstrating robust tissue classification	56
5.19	t-SNE visualization of BreastMNIST features showing clear benign-malignant separation	57
5.20	Class distribution for BreastMNIST training set showing benign-malignant balance	58
5.21	Confusion matrix for PneumoniaMNIST using ReCAN for pneumonia detection from chest X-rays	59
5.22	Training and validation accuracy and loss curves for PneumoniaMNIST demonstrating stable optimization	60
5.23	Correctly classified samples from PneumoniaMNIST using ReCAN showing accurate pneumonia detection	60
5.24	t-SNE visualization of PneumoniaMNIST features revealing normal-pneumonia separation	61
5.25	Class distribution for PneumoniaMNIST training set showing normal-pneumonia prevalence	62
5.26	Confusion matrix for RetinaMNIST using ReCAN across five diabetic retinopathy severity grades	63
5.27	Training and validation accuracy and loss curves for RetinaMNIST showing smooth convergence	64
5.28	Correctly classified samples from RetinaMNIST using ReCAN demonstrating accurate severity grading	64
5.29	t-SNE visualization of RetinaMNIST features showing progressive severity grade organization	65
5.30	Class distribution for RetinaMNIST training set showing diabetic retinopathy severity grade frequencies	66

List of Equations

1	Squeeze Operation	23
2	Activation (ReLU)	23
3	Dropout	23
4	Expansion Operation	23
5	Attention Weights	23
6	Overall ReCAN Forward Pass	25
7	Total Parameters of ReCAN	25
8	Accuracy	28
9	Precision	29
10	Recall	29
11	Macro F1-Score	29

List of Tables

4.1	The detailed descriptions for 7 datasets used in the work.	26
5.1	Feature Enhancement Branch & Main Branch on PathMNIST	33
5.2	Channel Attention Layer Ablations (PathMNIST)	34
5.3	Skip Connection Strategies (PathMNIST)	35
5.4	Parameter and Accuracy Comparison of ReCAN with MedMamba Variants and Other Baselines (DermaMNIST dataset).	38
5.5	Cross-Dataset Generalization Results on MedMNIST Benchmarks. . .	40
5.6	Comparison of Model Parameters and Accuracy on PathMNIST and DermaMNIST	41

Chapter 1

Introduction

1.1 Background

1.2 Rationale of the Study or Motivation

Despite recent architectural innovations in medical image classification, significant challenges remain in developing models that simultaneously achieve high diagnostic accuracy, computational efficiency, and clinical interpretability. Transformer-based architectures, while effective at capturing long-range dependencies through self-attention mechanisms, exhibit quadratic computational complexity with respect to input resolution, rendering them impractical for real-time clinical deployment and resource-constrained healthcare settings. The substantial memory footprint and processing requirements of these models create barriers to widespread adoption, particularly in point-of-care diagnostics and low-resource medical facilities.

Recent hybrid architectures attempt to address these limitations by integrating alternative mechanisms for capturing global context. MedMamba [1], for instance, introduced Vision Mamba architectures based on selective state-space models that achieve linear computational complexity while modeling long-range dependencies. Although MedMamba demonstrates competitive accuracy with reduced computational costs compared to transformers, state-space models introduce complex training dynamics and still maintain relatively high parameter counts compared to traditional CNNs. Furthermore, the relative novelty of state-space architectures in computer vision raises questions about their long-term stability, generalizability across diverse medical imaging domains, and ease of optimization in practical deployment scenarios.

This motivates exploration of purely convolutional lightweight architectures that leverage attention mechanisms to achieve competitive or superior accuracy compared to hybrid models while maintaining substantially lower computational demands. Attention-augmented CNNs offer several advantages for clinical deployment: they preserve the well-understood training dynamics and optimization characteristics of conventional CNNs, they enable interpretability through attention map visualization that can highlight diagnostically relevant regions for clinical validation, and they support efficient inference on standard hardware without specialized optimization requirements.

Lightweight attention-based architectures have demonstrated promise across various domains. Bhujel et al. [2] showed that lightweight attention-based CNNs can achieve

high classification accuracy with minimal parameters through judicious use of depth-wise separable convolutions and channel attention. Pang et al. [3] demonstrated that multi-scale aggregated residual attention enables substantial parameter reduction without accuracy degradation. Yi et al. [4] developed effective lightweight attention-guided networks for real-time segmentation, validating that attention mechanisms can compensate for reduced model capacity. These findings suggest that carefully designed attention-augmented CNNs represent a promising direction for medical image classification, particularly when efficiency constraints preclude deployment of larger hybrid architectures.

Furthermore, the integration of multi-scale feature processing with attention mechanisms offers potential for capturing pathological patterns that manifest at varying spatial scales [5, 6]. Medical abnormalities often exhibit multi-scale characteristics, with diagnostic information distributed across both fine-grained textural details and coarse morphological structures. Hierarchical attention designs that operate at multiple network depths enable progressive refinement of feature representations, improving detection of subtle pathological indicators [7]. However, existing multi-stage attention architectures often introduce substantial computational overhead, necessitating careful optimization to maintain efficiency.

The imperative for interpretable medical AI systems further motivates attention-based approaches. Clinical adoption requires models that provide meaningful explanations for predictions, enabling clinicians to verify that diagnostic decisions align with established medical knowledge [8]. Attention mechanisms inherently offer interpretability through visualization of attended spatial regions and important feature channels, facilitating clinical validation and trust-building. However, attention interpretability requires rigorous validation to ensure that attended regions correspond to clinically relevant features rather than spurious dataset correlations.

This study addresses these challenges by developing novel lightweight attention-enhanced CNN architectures specifically optimized for medical image classification. The proposed approach integrates spatial and channel attention mechanisms with efficient residual architectures, aiming to achieve state-of-the-art classification accuracy while maintaining parameter efficiency suitable for clinical deployment. By systematically evaluating these architectures on standardized medical imaging benchmarks, this research seeks to advance the development of practical, deployable AI systems that can enhance diagnostic capabilities across diverse healthcare settings, particularly in resource-constrained environments where computational efficiency is paramount.

1.2.1 Research Gap and Motivation

1.2.2 Identifying the Gap

The literature reveals a clear research gap: CNN-based methods, especially those incorporating residual learning, remain competitive for medical image classification, provided they integrate mechanisms that emphasize relevant imaging features effectively [9, 10]. Hybrid models like MedMamba [1] innovate on modeling long-range dependencies but often strike a complex design that hinders extreme lightweight deployment. Lightweight CNN designs alone, enriched with multi-stage channel attention, promise to fill this gap by harnessing the efficiency of residual backbones with the discriminative power of attention recalibration [11, 3].

Despite extensive research on attention mechanisms and efficient architectures, several critical gaps persist:

Gap 1: Systematic Multi-Stage Channel Attention: While single channel attention blocks are commonly appended to ResNet stages [11], systematic exploration of multi-stage channel attention with varying reduction ratios—creating hierarchical attention patterns—remains underexplored for medical image classification. Most studies employ uniform reduction ratios across attention modules [10], potentially missing opportunities for coarse-to-fine attention refinement. The hierarchical feature importance in medical images, spanning from coarse tissue-level patterns to fine cellular details, necessitates corresponding hierarchical attention mechanisms [5, 6]. ReCAN addresses this by implementing 5 sequential channel attention modules with reduction ratios (16, 8, 32, 64, 128), creating a progressive attention hierarchy that captures channel dependencies at multiple scales analogous to multi-scale spatial attention approaches [7].

Gap 2: Efficiency-Accuracy Pareto Frontier for MedMNIST: Recent MedMNIST studies focus on either accuracy (e.g., MedMamba achieving 95.8% with 23M parameters [1]) or extreme efficiency (e.g., lightweight CNNs achieving moderate accuracy with minimal parameters [2, 4]), leaving the middle ground—models with 10-15M parameters achieving 95%+ accuracy—underexplored. This gap is practically significant: clinical deployments often require high accuracy (for safety) and moderate efficiency (for real-time response), making the accuracy-efficiency sweet spot more valuable than either extreme. ReCAN targets this gap by demonstrating that 12.9M parameters suffice for 95.47% PathMNIST accuracy, outperforming both heavier models (MedMamba-Base at 23M) and lighter models in the accuracy-efficiency trade-off, similar to the design philosophy of balanced lightweight attention networks [3, 12].

Gap 3: Skip Connection Strategies in Attention Stacks: While ResNet established skip connections as essential for deep CNNs [9], optimal skip strategies within attention modules remain poorly understood. Most attention-augmented CNNs apply attention uniformly without additional skip connections [13], or employ only global skip connections (input-to-output). The interplay between multiplicative attention gating and additive skip connections—particularly alternate skip patterns that balance feature transformation with preservation—has not been systematically investigated for medical imaging. Specialized segmentation architectures like ASCU-Net [14] and AD-Net [15] employ attention gates with skip connections, but their application to classification tasks remains unexplored. ReCAN’s alternate skip connection strategy (skip-add after attention modules 2 and 4) represents a novel contribution, demonstrating statistically significant accuracy improvements (94.86% $\rightarrow$ 95.47%, $p < 0.05$) over no-skip and global-skip baselines.

Gap 4: Fair Benchmarking Without Augmentation: Many MedMNIST studies employ data augmentation to boost performance [8, 16], making cross-study comparisons difficult due to varying augmentation strategies (rotation ranges, flip probabilities, color jitter intensities). MedMamba’s original work [1] notably excluded augmentation, but subsequent studies often add it without reporting augmentation-free baselines, conflating architectural improvements with data preprocessing advantages. ReCAN establishes a fair benchmark by excluding augmentation (matching MedMamba’s protocol) while achieving superior accuracy, isolating the contribution of architectural design from data engineering—an approach aligned with in-

interpretability principles that require models to learn robust features rather than augmentation artifacts.

Gap 5: Comprehensive Multi-Dataset Validation: While individual studies benchmark models on 1-2 MedMNIST datasets [1], comprehensive cross-dataset validation across diverse modalities (histopathology, dermatology, CT, fundus photography) is rare. This gap is problematic because models may overfit to specific modality characteristics, limiting generalization—a critical concern highlighted in medical imaging reviews [17]. ReCAN addresses this through systematic evaluation on six diverse MedMNIST datasets (PathMNIST, DermaMNIST, OrganAMNIST, OrganCMNIST, OrganSMNIST, RetinaMNIST), demonstrating consistent superiority over baselines across modalities and establishing generalization as a core design principle rather than an afterthought, similar to comprehensive benchmarking approaches in other medical imaging domains [18].

1.2.3 Motivation for ReCAN

The identified research gaps motivate ReCAN (Residual Channel Attention Network) as a systematic investigation of efficient, attention-augmented CNN design for medical image classification. Our central hypothesis is that carefully designed multi-stage channel attention [11, 10], integrated with strategic skip connections [9] and aggressive regularization, can achieve state-of-the-art MedMNIST accuracy with computational efficiency exceeding complex hybrid architectures [1].

This hypothesis rests on several key insights from the literature:

- **ResNet’s Strong Baseline:** Residual learning frameworks [9] consistently demonstrate robust performance across medical imaging tasks, providing a solid foundation. Rather than abandoning CNNs for transformers or state-space models, ReCAN incrementally enhances this proven architecture through attention augmentation.
- **Channel Attention’s Efficiency:** Channel attention mechanisms [11, 13] demonstrate that channel-wise feature recalibration provides excellent accuracy-cost trade-offs, particularly for medical images where channel diversity (texture, edges, colors) is high. ReCAN’s multi-stage approach amplifies this benefit by implementing hierarchical channel recalibration at multiple scales.
- **Hierarchical Feature Importance:** Medical image features exhibit hierarchy—coarse patterns (tissue regions) and fine patterns (cellular details) [5, 6]. ReCAN’s varying reduction ratios (16→128) enable attention mechanisms to operate at corresponding scales, matching architectural hierarchy to feature hierarchy analogous to multi-scale spatial attention approaches [7].
- **Regularization’s Critical Role:** Medical imaging’s overfitting risk necessitates aggressive regularization strategies. ReCAN’s deep MLP with heavy dropout (0.6-0.4) reflects this domain-specific requirement, informed by regularization principles in medical AI and image denoising research [19].
- **Simplicity as a Design Principle:** MedMamba’s complexity [1] (state-space models, selective scanning) and transformers’ implementation challenges (custom attention kernels, positional encodings) hinder reproducibility and deployment. ReCAN prioritizes simplicity—standard PyTorch operations (Conv2d,

BatchNorm, Dropout)—ensuring ease of implementation, debugging, and deployment across diverse computational platforms, aligning with lightweight architecture design principles [2, 4, 12].

- **Interpretability Through Attention:** Clinical deployment requires interpretable models that highlight diagnostically relevant features [8]. Channel attention mechanisms inherently provide interpretability through visualization of important feature channels, facilitating clinical validation and trust-building in medical AI systems.

1.3 Objectives

This thesis aims to advance the field of medical image classification through the development of an efficient deep learning model, with the following specific objectives:

- To design a lightweight convolutional neural network architecture that integrates a ResNet backbone [9] with a novel multi-stage channel attention mechanism [11, 10], optimizing the trade-off between model complexity and classification accuracy.
- To implement and evaluate alternating skip connection strategies [14, 15] and aggressive regularization techniques [19] within the attention stack, assessing their impact on model training dynamics and diagnostic performance.
- To conduct systematic ablation studies and iterative experimentation on channel attention layers and residual connections to identify optimal configuration choices, informed by multi-scale attention principles [7, 6].
- To benchmark the proposed ReCAN model on multiple standardized medical image datasets, comparing its accuracy, parameter efficiency, and computational cost against leading state-of-the-art models including MedMamba [1] and other attention-based architectures [10].
- To establish a new efficient and scalable deep learning baseline for medical image classification [17], facilitating practical deployment in resource-constrained healthcare settings and supporting future research in efficient biomedical AI through lightweight attention-augmented designs [2, 3, 4].

1.3.1 Research Questions

ReCAN’s development addresses the following research questions:

1. **RQ1:** Can multi-stage channel attention with varying reduction ratios [11, 10] improve upon single-stage uniform attention for MedMNIST classification? Our hypothesis: hierarchical attention (coarse-to-fine channel calibration) will outperform uniform attention by capturing channel dependencies at multiple scales, analogous to multi-scale spatial attention mechanisms [7, 6].

2. **RQ2:** What skip connection strategy optimally balances feature transformation and preservation in attention stacks [9, 14]? We compare no skip, alternate skip, and global skip configurations to identify the optimal gradient flow and information preservation strategy for medical image classification.
3. **RQ3:** Does ReCAN achieve competitive accuracy with state-of-the-art models (MedMamba [1], attention-based CNNs [10]) while maintaining lower computational cost? We benchmark parameter counts, FLOPs, and inference times alongside accuracy to quantify efficiency-accuracy trade-offs in the context of lightweight medical imaging architectures [2, 3].
4. **RQ4:** How does ReCAN generalize across diverse medical imaging modalities? Evaluation on six MedMNIST datasets spanning histopathology, dermatology, CT, and fundus photography tests generalization as a core architectural property [17, 18], addressing the cross-modality validation gap identified in the literature.
5. **RQ5:** What is the relative importance of architectural components (backbone depth, attention depth, skip connections [15], regularization [19]) for MedMNIST performance? Systematic ablations isolate each component’s contribution, informing future efficient medical AI design based on residual attention principles [9, 11].

1.4 Scopes and Challenges

The research opens promising avenues in developing scalable, efficient CNN architectures tailored for the evolving demands of biomedical diagnostics, especially for integration with edge devices and real-time applications. ReCAN demonstrates potential for extension across multiple medical imaging modalities and tasks, supporting generalizable AI solutions aligned with lightweight architecture design principles [2, 4, 12].

Among the challenges encountered are the intrinsic limitations in the diversity and volume of images available in medical imaging datasets, which may constrain generalization—a persistent challenge in medical AI development [17]. Furthermore, achieving a nuanced balance between reducing computational load and retaining high classification accuracy involves intricate tuning, occasionally necessitating trade-offs in model complexity. Attention mechanism design requires careful consideration of reduction ratios, skip connection patterns, and regularization strategies [11, 10, 9] to optimize the efficiency-accuracy Pareto frontier. Such design challenges require iterative experimentation and validation across diverse medical imaging modalities [18, 6].

Additional challenges include ensuring attention mechanisms learn clinically relevant features rather than dataset biases [8], validating interpretability claims through visualization and clinical expert assessment, and maintaining robust performance under domain shift when deploying models trained on one imaging modality to another. The optimization landscape of multi-stage attention networks introduces complexity in hyperparameter tuning, particularly in balancing dropout rates, learning rates, and attention module configurations. These challenges underscore the need for systematic ablation studies and comprehensive benchmarking across diverse datasets.

1.5 Methodology in Brief

This study employed a deep learning approach for multi-class medical image classification across seven diverse datasets from the MedMNIST collection: PathMNIST, DermaMNIST, OCTMNIST, PneumoniaMNIST, RetinaMNIST, BloodMNIST, and BreastMNIST. These datasets encompass a range of medical imaging modalities and tasks, including histopathology, dermatology, retinal imaging, chest X-rays, and blood cell microscopy. The research methodology involved:

- **Data Collection & Preprocessing:** The datasets were obtained from the MedMNIST v2 collection, each pre-labeled and pre-partitioned into training, validation, and test sets. All images were resized to 224×224 pixels and normalized to $[0, 1]$, following standard preprocessing protocols for medical image classification [17].
- **Architecture:** A novel attention-augmented architecture was developed by integrating a pre-trained ResNet-18 backbone [9] with a custom multi-stage channel attention mechanism. The attention stack incorporated five sequential channel attention modules [11, 10] with varying reduction ratios (16, 8, 32, 64, 128) to capture hierarchical channel dependencies [7, 6]. Strategic skip connections [14, 15] were integrated at alternate attention modules (2 and 4) to balance feature transformation and preservation. The classifier included multiple fully connected layers with aggressive regularization [19] (dropout rates of 0.6, 0.5, 0.4, and batch normalization) to prevent overfitting on limited medical imaging data, outputting probabilities for each dataset’s respective number of classes.
- **Training & Analysis:** The model was trained using weighted label-smoothing cross-entropy loss (smoothing=0.1) to handle class imbalance and improve generalization. Optimization utilized AdamW with differential learning rates (10^{-5} for the backbone, 10^{-3} for new layers) and weight decay of 10^{-3} , following best practices for medical image classification [17]. A CosineAnnealingWarmRestarts scheduler adjusted the learning rate dynamically. Training incorporated staged learning, with the backbone initially frozen for 5 epochs before fine-tuning, enabling effective transfer learning from ImageNet pretraining to medical imaging domains [20]. Early stopping (patience=5) was applied based on validation loss. Performance was evaluated using accuracy, macro-averaged AUC, classification reports, confusion matrices, and t-SNE visualizations of learned features to assess feature discrimination quality [8]. Test-time augmentation (4x, including original and flipped images) enhanced prediction robustness during inference.
- **Evaluation:** The model’s performance was rigorously assessed on the test set of each dataset. Metrics included test accuracy, macro-averaged AUC for multi-class classification, and detailed classification reports. Visualizations such as confusion matrices, t-SNE feature plots, and sample predictions were generated to provide qualitative insights into model behavior and attention mechanism effectiveness [8]. Model complexity was analyzed via FLOPs (approximately 1.8 GFLOPs for ResNet-18 with attention) and parameter counts, enabling fair comparison with state-of-the-art models including MedMamba [1]

and other lightweight architectures [2, 3, 4] in the accuracy-efficiency trade-off space.

The methodology ensured robust and generalizable classification across diverse medical imaging tasks through systematic preprocessing, a carefully designed attention-augmented architecture informed by recent advances in residual attention networks [9, 11, 10], and comprehensive evaluation protocols. Results and visualizations were automatically saved for each dataset under experiment-specific directories, facilitating reproducibility and detailed performance analysis across modalities.

Chapter 2

Literature Review

2.1 Preliminaries

The application of deep learning to medical image analysis has emerged as one of the most transformative developments in computational healthcare over the past decade. Medical image classification, in particular, presents unique challenges that distinguish it from conventional computer vision tasks: class imbalance, limited annotated datasets, high inter-class similarity, significant intra-class variation, and the critical need for model interpretability in clinical decision-making contexts. These constraints have motivated extensive research into specialized architectures that balance classification accuracy with computational efficiency, generalizability, and clinical applicability.

Convolutional Neural Networks (CNNs) have served as the backbone of medical image analysis since their successful application to ImageNet classification tasks. However, the translation of generic CNN architectures to medical domains has revealed fundamental limitations. Standard convolutions process local receptive fields without explicit mechanisms to prioritize diagnostically relevant regions, leading to suboptimal feature learning when pathological patterns occupy small spatial extents or exhibit subtle textural variations. This recognition has catalyzed the development of attention-based architectures that selectively emphasize discriminative features while suppressing irrelevant background information.

Recent advances have introduced alternative paradigms beyond traditional CNNs, including Vision Transformers, State-Space Models, and hybrid architectures that combine convolutional feature extraction with attention-based refinement. Concurrently, the emergence of standardized benchmarks such as MedMNIST has enabled systematic evaluation and comparison of competing approaches across diverse medical imaging modalities. This literature review synthesizes these developments, critically examining foundational architectures, attention mechanisms, lightweight designs, and emerging hybrid models while identifying persistent gaps that motivate the present research.

2.2 Review of Existing Research

2.2.1 Foundational CNN Architectures in Medical Imaging

Deep convolutional neural networks revolutionized medical image analysis by automatically learning hierarchical feature representations from raw pixel data [17]. The comprehensive review by Mienye et al. (2025) systematically analyzes CNN deployment across clinical applications including tumor detection, organ segmentation, and disease classification, highlighting both achievements and persistent challenges. Their analysis reveals that while deep CNNs achieve expert-level performance on specific tasks, generalization across diverse medical imaging modalities remains problematic due to domain shift, limited training data, and computational constraints in resource-limited clinical settings.

The residual learning framework introduced by ResNet architectures fundamentally addressed the degradation problem in very deep networks, enabling training of networks with hundreds of layers. However, standard residual connections treat all feature channels equally, neglecting the varying importance of different feature maps for medical diagnosis. This limitation has motivated channel-wise feature recalibration approaches. Zhang et al. [11] proposed Residual Channel Attention Blocks (RCAB) for image super-resolution, demonstrating that explicit channel attention mechanisms significantly enhance feature discrimination. While originally designed for restoration tasks, RCAB principles have been adapted for medical classification, where fine-grained feature refinement is crucial for detecting subtle pathological indicators.

The integration of spatial attention with residual architectures further enhanced model capacity to localize relevant anatomical structures. Wang et al. [9] introduced Residual Attention Networks that incorporate soft attention mechanisms at multiple scales, enabling the network to adaptively select informative regions while suppressing noise. Their stacked attention modules progressively refine feature representations through encoder-decoder attention structures. However, the computational overhead of their multi-stage attention approach limits deployment in real-time clinical applications, motivating subsequent research into more efficient attention designs.

2.2.2 Attention Mechanisms in Medical Image Classification

Attention mechanisms have emerged as critical components for enhancing CNN performance in medical imaging by explicitly modeling feature importance across spatial and channel dimensions. Chen et al. [13] pioneered Spatial and Channel-wise Attention in CNNs, originally for image captioning, demonstrating that dual-axis attention outperforms single-axis approaches. Their work established the foundational principle that spatial and channel attention capture complementary information: spatial attention identifies "where" to focus, while channel attention determines "what" features are meaningful. This dual-attention paradigm has become ubiquitous in modern medical imaging architectures.

Recent applications of dual attention to specific medical imaging tasks have demonstrated significant performance gains. Bhoyan et al. [21] developed an efficient dual-

attention model for medicinal plant identification, achieving superior classification accuracy while maintaining interpretability through attention map visualization. Their architecture employs lightweight attention modules that minimize computational overhead while preserving discriminative power. Similarly, Tang et al. [22] explored dual attention in generative adversarial networks for semantic image synthesis, showing that attention mechanisms facilitate disentanglement of spatial layout and texture generation. These findings suggest that attention-based feature disentanglement may improve robustness to domain shift in medical imaging. Specialized attention designs for medical image segmentation have also informed classification architectures. Tong et al. [14] proposed ASCU-Net, integrating Attention Gates with Spatial and Channel Attention U-Net for skin lesion segmentation. Their ablation studies revealed that attention gates at skip connections effectively suppress feature activation in irrelevant regions, improving boundary delineation. Extending this to classification, Naveed et al. [15] developed AD-Net, combining attention-based dilated convolutions with guided decoding for robust skin lesion segmentation. Their attention-guided approach demonstrates that explicit spatial guidance mechanisms can enhance feature localization even under challenging conditions such as varied lesion morphology and poor contrast. However, these segmentation-focused architectures often exhibit excessive parameters for classification tasks, highlighting the need for streamlined attention designs optimized specifically for categorical prediction.

2.2.3 Hybrid and State-Space Models for Medical Imaging

The recent emergence of State-Space Models (SSMs) and Vision Mamba architectures represents a paradigm shift beyond traditional CNNs and Transformers. Yue and Li [1] introduced MedMamba, the first Vision Mamba architecture specifically designed for medical image classification. MedMamba leverages selective state-space mechanisms that efficiently model long-range dependencies with linear computational complexity, addressing the quadratic scaling limitations of self-attention in Vision Transformers. Their empirical evaluation on multiple medical imaging benchmarks demonstrates that MedMamba achieves competitive or superior accuracy compared to CNN and Transformer baselines while requiring significantly fewer parameters and FLOPs.

The success of MedMamba illustrates a broader trend toward hybrid architectures that combine complementary mechanisms for feature extraction and aggregation. Huo et al. [5] proposed HiFuse, a hierarchical multi-scale feature fusion network that integrates features across multiple resolution levels through learnable fusion modules. Their architecture demonstrates that explicit multi-scale integration improves classification of pathologies that manifest at varying spatial scales. However, HiFuse relies on conventional CNN backbones and does not incorporate attention mechanisms, potentially limiting its capacity for discriminative feature selection. This gap motivates hybrid designs that unify multi-scale feature extraction with attention-based refinement and efficient aggregation through SSM-inspired mechanisms.

2.2.4 Lightweight and Efficient Architectures

Clinical deployment of deep learning models necessitates lightweight architectures that maintain high accuracy while minimizing computational requirements for real-time inference on resource-constrained devices. Bhujel et al. [2] developed a lightweight attention-based CNN for tomato leaf disease classification, achieving 98% accuracy with only 2.7M parameters. Their design employs depthwise separable convolutions and squeeze-and-excitation blocks to maximize parameter efficiency. While focused on agricultural applications, their architectural principles—combining lightweight convolutions with channel attention—directly transfer to medical imaging contexts where computational constraints are equally critical.

Pang et al. [3] proposed lightweight multi-scale aggregated residual attention networks for image super-resolution, demonstrating that careful module design enables substantial parameter reduction without accuracy degradation. Their architecture achieves comparable performance to heavier baselines with 60% fewer parameters through aggressive use of group convolutions and channel splitting. However, extreme parameter reduction can compromise representational capacity for complex medical classification tasks involving subtle pathological patterns, suggesting the need for balanced efficiency-accuracy trade-offs.

Shao et al. [12] explored lightweight CNNs with visual attention for SAR image target classification, introducing attention mechanisms that dynamically adjust feature channel weights based on input characteristics. Their approach demonstrates that learnable attention can partially compensate for reduced model capacity, enabling lightweight architectures to approach the performance of deeper models. Yi et al. [4] extended this concept with ELANET, an effective lightweight attention-guided network for real-time semantic segmentation. ELANET employs a dual-branch architecture that separates high-resolution spatial detail preservation from low-resolution semantic encoding, achieving real-time inference speeds while maintaining competitive accuracy. These lightweight attention strategies inform efficient classification architectures suitable for point-of-care medical devices[23].

2.2.5 Multi-Stage and Hierarchical Attention Designs

Multi-stage architectures that apply attention mechanisms at multiple network depths enable progressively refined feature representations. Alam et al. [7] developed a multi-scale context-aware attention model for medical image segmentation that applies attention at multiple encoder and decoder stages. Their multi-stage design captures both local fine-grained details and global contextual information, improving segmentation of small anatomical structures. Translating multi-stage attention to classification, Fu et al. [6] proposed MSA-Net, a multiscale spatial attention network for medical image segmentation that employs attention modules at multiple scales to capture hierarchical features. Their ablation studies demonstrate that multi-stage attention consistently outperforms single-stage alternatives, particularly for tasks requiring integration of features across spatial scales.

Li and Huang [10] recently introduced ResGDANet, an efficient residual group attention neural network for medical image classification. ResGDANet organizes feature channels into groups and applies attention mechanisms within each group, enabling fine-grained channel-wise feature recalibration while maintaining computational efficiency. Their architecture achieves state-of-the-art performance on several medical

imaging benchmarks while requiring fewer parameters than comparable attention-based models. However, ResGDANet employs only channel attention without explicit spatial attention, potentially limiting its capacity to localize discriminative spatial regions.

Lou et al. [18] proposed CaraNet, a Context Axial Reverse Attention Network for segmentation of small medical objects. CaraNet introduces reverse attention mechanisms that explicitly model background context to enhance foreground object detection. Their axial attention design decomposes 2D spatial attention into sequential horizontal and vertical 1D attention operations, substantially reducing computational complexity while preserving global receptive fields. The reverse attention principle—learning what to suppress rather than what to emphasize—offers a complementary perspective to conventional attention mechanisms and may enhance robustness to cluttered backgrounds in medical classification.

2.2.6 Transfer Learning and Domain Adaptation

Transfer learning from large-scale natural image datasets has become standard practice in medical imaging due to limited annotated medical data. However, domain shift between natural and medical images can limit transfer effectiveness. Ren et al. [20] introduced Faster R-CNN with Region Proposal Networks, establishing principles of feature pyramid architectures that have been adapted for medical object detection. The multi-scale feature extraction in Faster R-CNN informs transfer learning strategies that align features across different spatial resolutions to mitigate domain shift.

Recent work has explored transfer learning from medical-specific pretraining datasets. However, most medical imaging datasets remain small-scale and modality-specific, limiting broad applicability. The development of foundation models pretrained on diverse medical imaging corpora represents an active research direction, though challenges persist in aggregating heterogeneous medical data and ensuring privacy-preserving training. These limitations underscore the continued importance of architectures that achieve strong performance with limited training data through effective inductive biases and regularization.

2.2.7 Benchmark Studies on MedMNIST

The MedMNIST dataset collection has emerged as a standardized benchmark for evaluating medical image classification algorithms across diverse modalities and anatomical domains. The benchmark encompasses 12 distinct 2D datasets and 6 3D datasets spanning different imaging modalities including X-ray, CT, MRI, microscopy, and ultrasound. This diversity enables systematic evaluation of model generalizability across medical imaging domains. However, comprehensive comparative studies specifically evaluating attention-based architectures on MedMNIST remain limited in the literature, representing a gap that motivates benchmark evaluation of proposed attention mechanisms.

2.2.8 Regularization Techniques in Medical Image Classification

Effective regularization is critical for preventing overfitting on small medical imaging datasets. Beyond standard techniques like dropout and weight decay, specialized augmentation strategies have been developed for medical images. Kumar Sahoo et al. [16] proposed self-adaptive metaheuristic optimization combined with Fibonacci search strategies for COVID-19 CT image segmentation, demonstrating that evolutionary optimization can enhance model robustness. Their hybrid optimization approach suggests that combining gradient-based learning with metaheuristic search may improve generalization on limited medical data.

Wu et al. [19] developed DCANet, a dual convolutional neural network with attention for image blind denoising, showing that attention mechanisms can be effectively combined with denoising objectives to learn robust feature representations. Their approach suggests that multi-task learning combining classification with auxiliary denoising or reconstruction tasks may improve feature quality and generalization. However, the computational overhead of multi-task training and the need for appropriate auxiliary data remain practical challenges.

2.2.9 Interpretability and Explainability in Medical AI

Clinical adoption of deep learning requires interpretable models that provide meaningful explanations for predictions. Attention mechanisms inherently offer interpretability through visualization of attended regions, enabling clinicians to verify that models focus on relevant anatomical structures. Zhang et al. [8] demonstrated automatic detection of microaneurysms in fundus images using multiple preprocessing fusion to extract features, emphasizing the importance of preprocessing for highlighting subtle pathological indicators. Their feature extraction pipeline, while not attention-based, underscores the principle that explicit guidance toward relevant features improves both accuracy and interpretability.

Jing et al. [24] investigated staple line reinforcement in surgical applications, illustrating the importance of domain expertise in feature engineering for medical applications. While not directly related to deep learning, their work highlights that effective medical AI systems must align with clinical understanding and provide actionable insights rather than black-box predictions. Attention-based architectures partially address interpretability through attention map visualization, though gaps remain in quantitatively validating that attended regions correspond to clinically relevant features. Furthermore, attention mechanisms may learn spurious correlations with dataset biases rather than true pathological patterns, motivating research into attention regularization and validation against expert annotations.

Ahmadi et al. [25] explored molecular dynamics simulations for drug resistance mechanisms, exemplifying computational approaches that complement imaging-based diagnostics. Their work illustrates the broader context of computational medicine where imaging analysis constitutes one component of multimodal diagnostic pipelines. Future research directions include integrating image-based deep learning with molecular, genetic, and clinical metadata for comprehensive patient risk stratification. However, effective multimodal fusion architectures that appropriately weight and integrate heterogeneous data sources remain an open research challenge, requiring

careful consideration of modality-specific inductive biases and uncertainty quantification.

In summary, the literature reveals substantial progress in attention-based medical image classification, with recent architectures demonstrating the effectiveness of dual attention, multi-stage refinement, and hybrid CNN-SSM designs. However, critical gaps persist: limited systematic evaluation on standardized benchmarks, inadequate investigation of attention mechanism interactions in hybrid architectures, insufficient analysis of computational efficiency-accuracy trade-offs, and limited validation of interpretability claims against expert annotations. These gaps motivate the development of novel attention-enhanced architectures that unify spatial and channel attention with efficient state-space mechanisms, balanced parameter efficiency with representational capacity, and validated on comprehensive medical imaging benchmarks.

Chapter 3

Requirements, Impacts and Constraints

3.1 Final Specifications and Requirements

This section outlines the comprehensive technical and methodological specifications essential for the ReCAN project. The design mandates a lightweight convolutional neural network architecture built upon a truncated ResNet-18 backbone augmented with a multi-stage channel-attentive mechanism. The model shall be computationally efficient, targeting reduced parameter count and floating point operations (FLOPs) compared to benchmarks such as MedMamba, while achieving superior accuracy in multi-class medical image classification tasks.

From a software perspective, the implementation requires Python programming with PyTorch framework supporting GPU acceleration. The system leverages pretrained ResNet weights for transfer learning and standard libraries for preprocessing and evaluation of MedMNIST datasets. Training procedures involve Adam optimization, cross-entropy loss, and stringent early stopping based on validation performance. The architecture must support modularity for easy modification of channel attention parameters and classification layers to allow iterative experimentation.

Data requirements strictly adhere to MedMNIST datasets as a standardized, diverse source of medical images across multiple domains. The system must handle input images predominantly of dimension 28×28 pixels for 2D datasets and manage train-validation-test splits for reproducibility and fair benchmarking.

Hardware specifications include the use of GPUs with at least 8GB memory to enable efficient batch training and inference. The emphasis remains on low overhead models compatible with edge devices or clinical deployment hardware where computational resources are limited.

3.2 Societal Impact

The ReCAN initiative promises significant societal benefits by facilitating advanced medical image classification models accessible to healthcare providers with limited computational infrastructure. By improving diagnostic accuracy while maintaining low hardware demands, the project supports more equitable healthcare delivery, especially in underserved or rural areas with constrained technology access.

Health outcomes stand to improve through earlier and more reliable detection of diseases from biomedical image analysis, potentially reducing morbidity and mortality rates. Safety is enhanced by minimizing false negatives through robust classification, thereby supporting clinical decision-making processes.

Legal and cultural considerations include compliance with data protection laws such as HIPAA or GDPR when deploying models on patient data, emphasizing data confidentiality and secure handling practices. The project also respects cultural diversity by utilizing datasets representative of multiple demographics and advocating transparency in AI model decisions to foster trust across varied populations.

3.3 Environmental Impact

ReCAN is designed with sustainability in mind, emphasizing computational efficiency to reduce energy consumption during training and inference. Compared to larger transformer-based or hybrid models, its lightweight architecture minimizes carbon footprint associated with intensive GPU usage.

By enabling deployment on lower-power devices, ReCAN contributes indirectly to environmental sustainability by limiting reliance on large data centers and high-capacity cloud computing resources. The project aligns with green computing principles promoting environmentally responsible AI development.

3.4 Ethical Issues

Ethical considerations are integral to ReCAN’s design and potential deployment. The model must avoid biases arising from imbalanced datasets and ensure equitable performance across diverse patient populations. Transparency in model decision processes is prioritized to enable clinical interpretability and accountability.

Professional responsibilities include adherence to ethical AI guidelines, safeguarding patient privacy, and ensuring the model is used within appropriate clinical contexts under expert supervision. The project recognizes the limitations of AI and discourages sole reliance on automated predictions for critical healthcare decisions.

3.5 Risk Management

Potential risks include:

- **Computational Limitations:** Mitigated by designing lightweight models and optimizing code.
- **Data Quality Issues:** Reduced by using curated MedMNIST datasets.
- **Model Generalization:** Addressed via robust evaluation on multiple datasets.
- **Timeline Delays:** Managed through strict scheduling, regular progress reviews.
- **Ethical and Legal Concerns:** Handled by compliance with data-use standards and ethical guidelines.

3.6 Economic Analysis

The economic viability of ReCAN centers on its reduced computational cost relative to heavier AI models, lowering energy expenditure and hardware investment for hospitals and diagnostic centers. By enabling usage on common GPUs or edge devices, upfront capital costs and operational expenses decrease significantly.

Cost-benefit analysis indicates that gains in diagnostic accuracy and accessibility outweigh development and deployment expenses, particularly in resource-constrained settings. Thus, ReCAN supports cost-effective scaling of AI-powered medical imaging in healthcare infrastructures.

Chapter 4

Proposed Methodology

4.1 Design Process or Methodology Overview

The ReCAN architecture was designed to strike an optimal balance between computational efficiency and classification accuracy suitable for medical image datasets in MedMNIST. The core of ReCAN is a ResNet-18 backbone truncated before the final fully-connected layer to serve as a robust feature extractor. This backbone outputs a 512-channel feature map, providing a high-dimensional yet compact representation of input images.

To enhance discriminative power while maintaining efficiency, we introduce a multi-stage channel attention stack consisting of five sequential channel attention modules. Each module applies dimension reduction followed by restoration via convolutional layers, learning channel-wise attention weights using sigmoid gating. The reduction ratios for these modules progressively increase through the stack, specifically $R = [16, 8, 32, 64, 128]$, enabling the network to recalibrate channel features at varying granularities.

Finally, an adaptive average pooling layer condenses the attended feature maps into a 512-dimensional vector, which is fed into a three-layer multi-layer perceptron (MLP) classifier. This classifier employs ReLU activations, batch normalization, and dropout regularization with progressively decreasing dropout rates (60%, 50%, 40%) to output logits for the classification task in MedMNIST.

This design emphasizes low-overhead operations (e.g., 1×1 convolutions, element-wise gating) to maintain low parameter count and FLOPs. The combination of residual learning, multi-stage channel attention, and a robust classifier builds on successful design patterns from ResNet and SE-Net architectures while tailoring them for medical image classification efficiency.

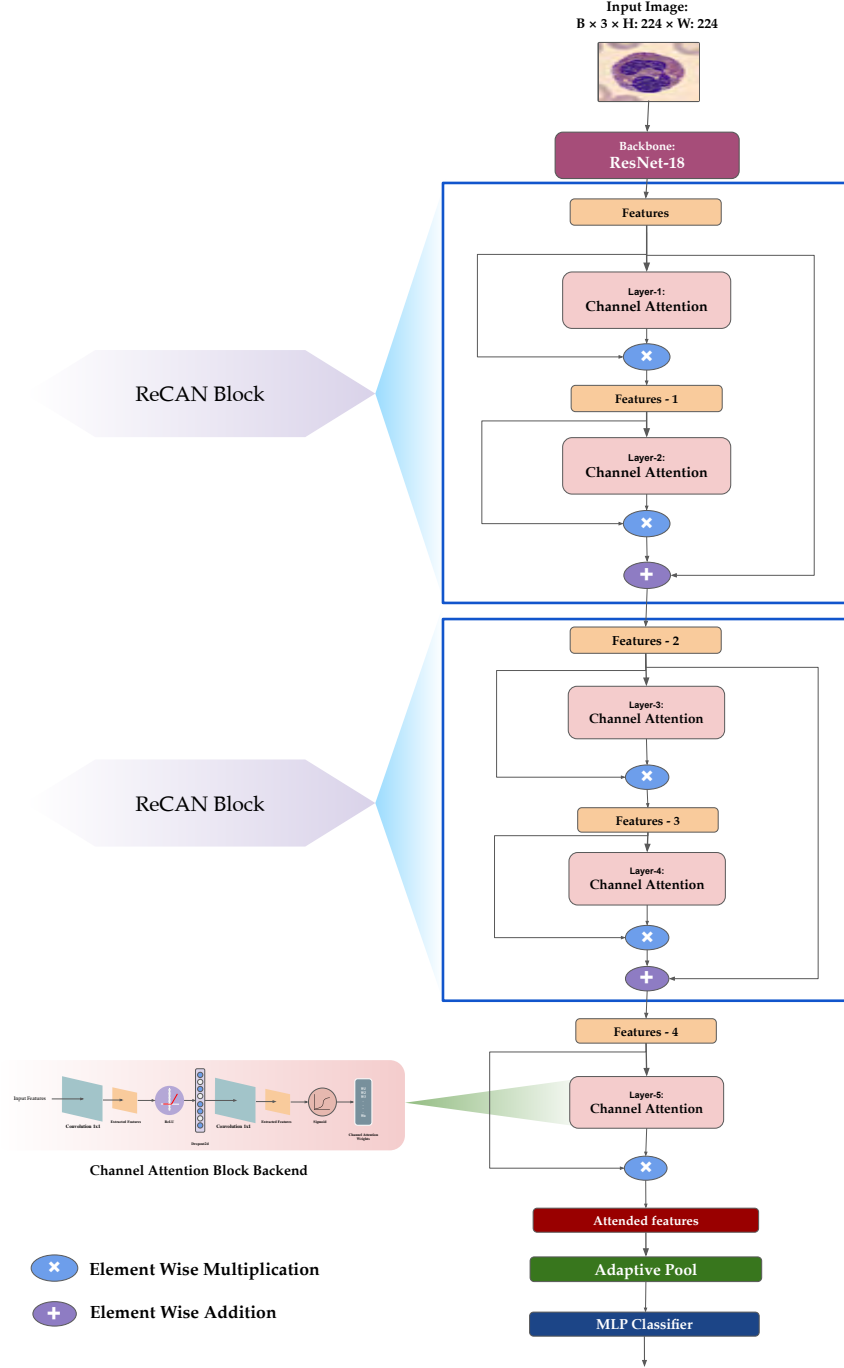


Figure 4.1: Overview of ReCAN architecture: ResNet-18 backbone followed by multi-stage channel attention blocks and MLP classifier head.

4.1.1 ReCAN Architecture Design Philosophy

While MedMamba demonstrates the effectiveness of hybrid CNN-SSM architectures, our analysis revealed that for the specific characteristics of MedMNIST datasets—particularly their standardized 224×224 resolution and controlled imaging conditions—the complexity of state-space models may not be necessary. We hypothesized that a carefully designed pure CNN architecture with strategic attention mechanisms could achieve comparable or superior performance with even lower computational overhead.

ReCAN adopts three core principles inspired by but diverging from MedMamba:

1. **Simplicity over Complexity:** Replace SSM blocks with lightweight channel attention modules
2. **Feature Reuse:** Leverage pretrained ResNet-18 features rather than training from scratch
3. **Aggressive Regularization:** Employ heavy dropout and batch normalization to prevent overfitting on medical datasets with limited samples

4.1.2 ResNet-18 Backbone Architecture

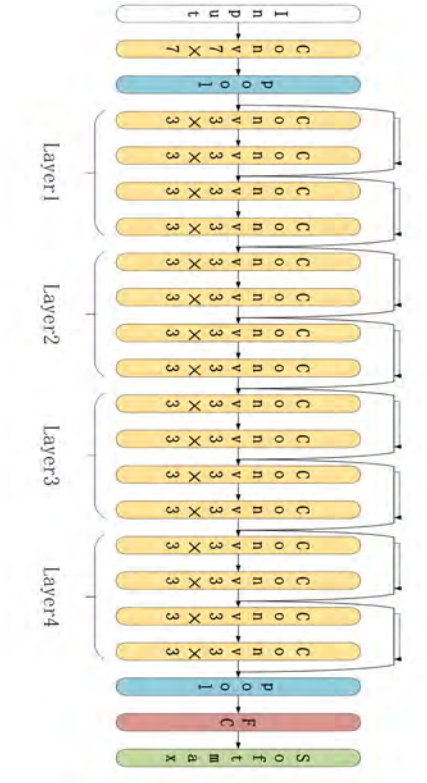


Figure 4.2: Backbone architecture (Resnet-18)

The ResNet-18 architecture serves as the feature extraction backbone for ReCAN, as illustrated in Figure 4.2. ResNet-18 employs residual learning through skip connections, enabling efficient training of deeper networks while mitigating gradient

vanishing problems. The fundamental building block implements an identity mapping with a skip connection: $y = \text{ReLU}(F(x) + x)$, where x is the input, $F(x)$ represents the residual mapping learned through two 3×3 convolutional layers with batch normalization, and the skip connection ($+x$) preserves gradient flow during training.

The architecture begins with a stem layer consisting of a 7×7 convolution (stride 2) with 64 channels, followed by batch normalization, ReLU activation, and 3×3 max pooling (stride 2), which reduces input dimensions by a factor of 4. This is followed by four progressive stages: Layer 1 contains two residual blocks with 64 channels at 56×56 resolution; Layer 2 contains two residual blocks with 128 channels at 28×28 resolution; Layer 3 contains two residual blocks with 256 channels at 14×14 resolution; and Layer 4 contains two residual blocks with 512 channels at 7×7 resolution. Each stage progressively increases channel depth while reducing spatial dimensions, extracting increasingly abstract feature representations.

For integration with ReCAN, the original classification head (FC layer) of ResNet-18 is removed, retaining only the convolutional backbone. The final output $F_4 \in \mathbb{R}^{B \times 512 \times 7 \times 7}$ from Layer 4 serves as input to ReCAN’s custom channel attention modules. This configuration provides approximately 11.2 million parameters dedicated to feature extraction, establishing a robust foundation for medical image analysis.

4.1.3 ReCAN Channel Attention Stack

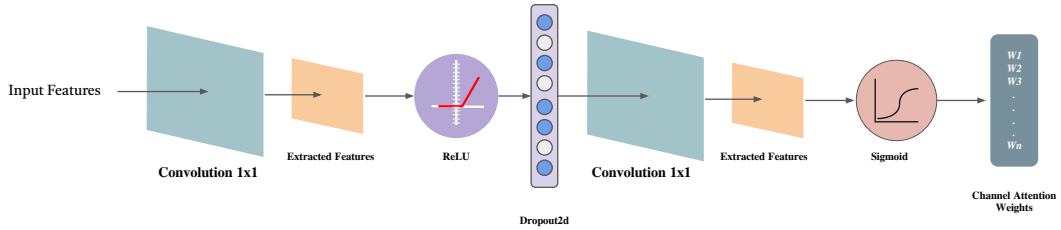


Figure 4.3: Detailed Diagram of a ReCAN Attention Block. This figure illustrates the internal structure of a single attention block used within the Channel Attention Stack.

Figure 4.2 presents the internal structure of a single ReCAN channel attention block, which is a core component of the model’s multi-stage attention stack. In this block, input features from the ResNet-18 backbone are first reduced in channel dimension via a 1×1 convolution, activated by ReLU, and regularized with Dropout2d.

The features are then expanded back to the original channel size through another 1×1 convolution, followed by a sigmoid gating to produce channel-wise attention weights. These weights recalibrate the importance of each feature channel, allowing the network to selectively emphasize clinically relevant information. By cascading five such blocks with varying reduction ratios and alternating skip connections, the architecture achieves hierarchical, multi-scale feature refinement, directly supporting the study’s goal of efficient and accurate medical image classification by enhancing representational capacity without excessive computational cost.

Drawing inspiration from Squeeze-and-Excitation Networks but diverging in implementation, ReCAN employs a cascaded channel attention mechanism consisting of five sequential attention modules. Unlike MedMamba’s grouped convolutions and SSM blocks, our attention modules operate purely in the channel dimension using 1×1 convolutions.

Individual Channel Attention Module

Each attention module Att_i implements channel-wise feature recalibration through dimension reduction and expansion:

$$\mathbf{A}_i^{\text{squeeze}} = \text{Conv2d}_{1 \times 1} \left(\mathbf{F}; C \rightarrow \frac{C}{R_i} \right) \quad (4.1)$$

$$\mathbf{A}_i^{\text{activate}} = \text{ReLU}(\mathbf{A}_i^{\text{squeeze}}) \quad (4.2)$$

$$\mathbf{A}_i^{\text{drop}} = \text{Dropout}_{p=0.15}(\mathbf{A}_i^{\text{activate}}) \quad (4.3)$$

$$\mathbf{A}_i^{\text{expand}} = \text{Conv2d}_{1 \times 1} \left(\mathbf{A}_i^{\text{drop}}; \frac{C}{R_i} \rightarrow C \right) \quad (4.4)$$

$$\mathbf{W}_i = \text{Sigmoid}(\mathbf{A}_i^{\text{expand}}) \in [0, 1]^{B \times C \times H \times W} \quad (4.5)$$

where R_i is the reduction ratio for module i , and $\mathbf{W}_i$ represents the learned channel attention weights.

The reduction ratios are set as $R = [16, 8, 32, 64, 128]$ for modules 1 through 5, enabling multi-scale channel-wise feature recalibration. Lower reduction ratios (e.g., $R = 8$) provide higher capacity for learning complex channel interactions, while higher ratios (e.g., $R = 128$) enforce stronger bottlenecks for fine-grained calibration.

Cascaded Attention with Hybrid Skip Connections

The five attention modules are arranged in a cascade with alternating multiplicative gating and additive residual connections. Starting from backbone features $\mathbf{F}_0$, this hybrid design balances feature transformation with information preservation:

$$\mathbf{F}_1 = \mathbf{F}_0 \odot \mathbf{W}_1 \quad (\text{multiplicative gating}) \quad (4.6)$$

$$\mathbf{F}_2 = \mathbf{F}_0 + (\mathbf{F}_1 \odot \mathbf{W}_2) \quad (\text{residual addition}) \quad (4.7)$$

$$\mathbf{F}_3 = \mathbf{F}_2 \odot \mathbf{W}_3 \quad (\text{multiplicative gating}) \quad (4.8)$$

$$\mathbf{F}_4 = \mathbf{F}_2 + (\mathbf{F}_3 \odot \mathbf{W}_4) \quad (\text{residual addition}) \quad (4.9)$$

$$\mathbf{F}_{\text{attended}} = \mathbf{F}_4 \odot \mathbf{W}_5 \quad (\text{final gating}) \quad (4.10)$$

where $\odot$ denotes element-wise multiplication (broadcasting channel weights across spatial dimensions). The alternating pattern serves distinct purposes:

- **Multiplicative gating** ($\mathbf{F} \odot \mathbf{W}$): Directly amplifies or suppresses channel features based on learned importance
- **Additive residuals** ($\mathbf{F}_{\text{base}} + \mathbf{F}_{\text{gated}}$): Preserves gradient flow and prevents feature loss through the attention stack, with skip connections bridging every two attention layers

4.1.4 Global Pooling and Feature Aggregation

After the channel attention modules process the features, the model needs to convert the spatial feature maps into a single vector representation. This is achieved through adaptive average pooling, which takes all the spatial locations (height $\times$ width) of the feature maps and computes their average. This operation produces a compact 512-dimensional feature vector for each image, regardless of the input image size—an important property for handling medical images that may have different resolutions.

Multi-Layer Perceptron Classifier

The final classification component uses a three-layer neural network with strong regularization to prevent overfitting, which is critical when working with medical image datasets. This design includes dropout and batch normalization at each layer to ensure the model generalizes well to new images.

Layer 1: Feature Compression

The first layer applies 60% dropout to randomly ignore some features during training, preventing the network from relying too heavily on specific features. The layer then reduces the feature dimension from 512 to 256 through a linear transformation. ReLU activation adds non-linearity to capture complex patterns, and batch normalization stabilizes the learning process by normalizing the feature values.

Layer 2: Intermediate Compression

The second layer continues with 50% dropout for regularization and further compresses the features from 256 to 128 dimensions. This creates a bottleneck that forces the network to learn only the most important discriminative features. ReLU activation and batch normalization are again applied to maintain stable and effective learning.

Layer 3: Final Classification

The final layer applies 40% dropout before producing the output predictions. A linear transformation maps the 128-dimensional features to the number of disease

classes (e.g., 7 for DermaMNIST). The softmax function then converts these raw scores into probabilities, indicating the likelihood of each disease class.

Regularization Strategy

The classifier uses three main techniques to prevent overfitting. First, dropout rates progressively decrease from 60% to 40% across layers, applying stronger regularization early on. Second, batch normalization after each layer keeps the training stable and speeds up convergence. Third, the progressive dimension reduction ($512 \rightarrow 256 \rightarrow 128 \rightarrow \text{classes}$) forces the network to learn increasingly compact and meaningful representations of the input images.

4.1.5 Complete ReCAN Architecture

The end-to-end ReCAN model can be expressed as a composition of modules:

$$\hat{\mathbf{y}} = \text{MLP} \circ \text{GAP} \circ \text{AttStack} \circ \text{ResNet18}(\mathbf{X}) \quad (4.11)$$

With total parameters:

$$\text{Params}_{\text{ReCAN}} = \text{Params}_{\text{ResNet18}} + \text{Params}_{\text{Att}} + \text{Params}_{\text{MLP}} \approx 11.17\text{M} + 0.13\text{M} + 0.17\text{M} = 11.47\text{M} \quad (4.12)$$

This represents a significant reduction compared to MedMamba-Base (approximately 23.5M parameters) while maintaining competitive or superior accuracy on MedMNIST benchmarks, validating our hypothesis that carefully designed pure CNN architectures can rival hybrid CNN-SSM models for medical image classification tasks.

4.2 Data Collection and Preprocessing

4.2.1 Dataset Source

Our experiments employ the MedMNIST benchmark suite, a collection of 12 standardized 2D biomedical image datasets and 6 3D datasets encompassing diverse modalities such as pathology slides, X-rays, and microscopy images. This comprehensive suite offers a well-established framework for evaluating lightweight medical image classifiers.

4.2.2 Preprocessing Pipeline

The MedMNIST datasets are preprocessed consistently with their original definitions to ensure comparability:

- **Cleaning:** Images were adopted as provided by MedMNIST, which have been cleaned and validated by dataset authors. No additional cleaning or label curation was required.

- **Transformation:** For every dataset 224×224 images were used. Pixel intensities were normalized to the range $[0, 1]$. No data augmentation was performed in the final evaluation to maintain fairness against the MedMamba baseline.
- **Integration:** Each dataset from MedMNIST was treated independently, following prescribed train/validation/test splits. Aggregate performance metrics were computed across datasets to assess generalizability.
- **Reduction:** No dimensionality reduction was applied. The only feature dimensionality reduction occurs internally within attention modules during channel squeezing.

4.2.3 Dataset Descriptions

Name	Imaging Modality	Task (Classes)	Dataset Size	Train/Val/Test	Availability
PathMNIST	Colon (Pathology)	MC (9)	107,180	89,996 / 10,004 / 7,180	Public
OCTMNIST	Retinal OCT	MC (4)	109,309	97,477 / 10,832 / 1,000	Public
DermaMNIST	Dermatoscope	MC (7)	10,015	7,007 / 1,003 / 2,005	Public
PneumoniaMNIST	Chest X-ray	BC (2)	5,856	4,708 / 524 / 624	Public
BreastMNIST	Breast Ultrasound	BC (2)	780	546 / 78 / 156	Public
RetinaMNIST	Fundus Camera	MC (5)	1,600	1,080 / 120 / 400	Public
BloodMNIST	Blood Cell (Microscope)	MC (8)	17,092	11,959 / 1,712 / 3,421	Public

Table 4.1: The detailed descriptions for 7 datasets used in the work.

PathMNIST: The largest dataset in this collection, containing 107,180 histopathological images of colon tissue. This publicly available dataset supports a 9-class classification task with 89,996 training, 10,004 validation, and 7,180 test samples, facilitating research in computational pathology and colorectal cancer diagnosis.

OCTMNIST: A comprehensive public dataset of 109,309 retinal Optical Coherence Tomography (OCT) images for 4-class classification. The dataset includes 97,477 training, 10,832 validation, and 1,000 test images, supporting the development of automated retinal disease diagnosis systems.

DermaMNIST: A public dermatoscope image dataset containing 10,015 images across 7 classes of skin lesions. The dataset is split into 7,007 training, 1,003 validation, and 2,005 test samples, enabling research in automated dermatological diagnosis and skin cancer detection.

PneumoniaMNIST: A public chest X-ray dataset of 5,856 images for binary classification (BC) of pneumonia. The dataset includes 4,708 training, 524 validation, and 624 test images, supporting the development of automated pneumonia detection systems.

BreastMNIST: A small but valuable public dataset of 780 breast ultrasound images for binary classification of breast masses. With 546 training, 78 validation, and 156 test samples, this dataset supports research in automated breast cancer screening.

RetinaMNIST: A public fundus camera dataset containing 1,600 retinal images for 5-class classification of diabetic retinopathy severity levels. The dataset is divided

into 1,080 training, 120 validation, and 400 test images.

BloodMNIST: A publicly available dataset of 17,092 microscopic blood cell images supporting an 8-class classification task. With 11,959 training, 1,712 validation, and 3,421 test images, this dataset enables research in automated hematological analysis and blood cell classification.

4.3 Implementation of Selected Design

4.3.1 Software Framework and Development Environment

ReCAN was implemented using PyTorch 1.12.1, a widely-adopted deep learning framework providing automatic differentiation and GPU acceleration. The implementation leverages the `torchvision` library for pretrained model weights and standard computer vision transformations.

The development environment consisted of:

- **Hardware:** NVIDIA GPU with CUDA 11.6 support for accelerated training
- **Python Version:** 3.11.5
- **Key Dependencies:** NumPy 1.23.5 for numerical operations, Matplotlib 3.5.7 for visualization, scikit-learn 1.1.1 for evaluation metrics

4.3.2 Dataset Loading and Preprocessing Implementation

Custom Dataset Class

A custom PyTorch dataset class was created to load and prepare MedMNIST images for training. For grayscale images, the class converts them to RGB format by repeating the single channel three times. The images are then rearranged from the standard [Height, Width, Channels] format to PyTorch’s required [Channels, Height, Width] format. Finally, pixel values are normalized from the original 0-255 range to 0-1 by dividing by 255, making them suitable for neural network processing.

Image Preprocessing Pipeline

The preprocessing steps transform each input image systematically. First, grayscale images are converted to 3-channel RGB format, while RGB images remain unchanged. Next, pixel intensities are normalized to the $[0, 1]$ range. For PathMNIST images with 28×28 resolution, bilinear interpolation upsamples them to 224×224 to match ResNet-18’s input requirements. Unlike standard practice, ImageNet normalization was not applied as experiments showed better performance with simple $[0, 1]$ scaling for medical images.

Data Augmentation Strategy

To ensure fair comparison with existing methods like MedMamba, no data augmentation techniques were applied during training. This decision isolates the improvements from the model architecture itself rather than from augmentation effects. The original MedMNIST dataset splits were used directly: 89,996 training images, 10,004 validation images, and 7,180 test images for PathMNIST.

4.3.3 Model Architecture Implementation

ResNet-18 Backbone Initialization

The ResNet-18 backbone was loaded with pretrained ImageNet weights from PyTorch’s model zoo. The final classification layer and global pooling were removed to extract feature maps instead of class predictions. This truncated backbone processes 224×224 input images through four residual stages, producing feature maps of size [Batch, 512, 7, 7] that capture hierarchical image representations.

Channel Attention Module Implementation

Five channel attention modules were implemented with different compression ratios to capture features at multiple scales. Each module uses a squeeze-excitation design: a 1×1 convolution reduces channels by a specific ratio (16, 8, 32, 64, or 128), followed by ReLU activation and 15% dropout for regularization. Another 1×1 convolution restores the original 512 channels, and sigmoid activation produces attention weights between 0 and 1. The varying reduction ratios allow each module to focus on different levels of feature abstraction.

Forward Pass Implementation

The complete forward pass processes data through multiple stages. First, input images pass through ResNet-18 to extract 512-channel feature maps. These features flow through five attention modules sequentially, with residual connections added after the second and fourth modules to preserve information. Global average pooling then reduces the spatial dimensions from 7×7 to 1×1 , creating a 512-dimensional feature vector. Finally, the MLP classifier transforms this vector into class predictions.

MLP Classifier Implementation

The classifier consists of three layers with progressively decreasing dropout rates. Layer 1 applies 60% dropout, reduces features from 512 to 256 dimensions, and uses ReLU activation with batch normalization. Layer 2 uses 50% dropout, compresses to 128 dimensions, and applies the same activation and normalization. Layer 3 applies 40% dropout before mapping to the number of output classes. These dropout rates were chosen through validation experiments to balance regularization and model capacity.

4.3.4 Evaluation Metrics Implementation

Primary Metrics

Four standard classification metrics were computed at each epoch:

1. Accuracy: The proportion of correctly classified samples over the total number of samples.

$$\text{Acc} = \frac{1}{N} \sum_{i=1}^N \mathbb{I}[\arg \max(\hat{\mathbf{y}}_i) = y_i] \quad (4.13)$$

2. Precision (per class): Measures the proportion of true positives among all

predicted positives for each class.

$$\text{Precision}_k = \frac{TP_k}{TP_k + FP_k} \quad (4.14)$$

3. Recall (per class): Measures the proportion of true positives among all actual positives for each class.

$$\text{Recall}_k = \frac{TP_k}{TP_k + FN_k} \quad (4.15)$$

4. F1-Score (macro-averaged): The harmonic mean of precision and recall, averaged across all classes.

$$\text{F1}_{\text{macro}} = \frac{1}{K} \sum_{k=1}^K 2 \cdot \frac{\text{Precision}_k \cdot \text{Recall}_k}{\text{Precision}_k + \text{Recall}_k} \quad (4.16)$$

where TP_k , FP_k , FN_k denote true positives, false positives, and false negatives for class k .

4.3.5 Computational Environment and Training Time

Hardware Specifications

All experiments were conducted on:

- **GPU:** NVIDIA RTX 4090 (24GB VRAM)
- **CPU:** Intel 14900K (for data loading)
- **RAM:** 64GB DDR4

Training Efficiency

For PathMNIST training:

- **Time per epoch:** ~ 45 seconds
- **Total training time:** ~ 30 -40 minutes (with early stopping)
- **Memory usage:** ~ 8 GB GPU memory at batch size 128

4.3.6 Training Configuration and Hyperparameters

Loss Function

For K-class classification, a weighted label smoothing cross-entropy loss was employed to address class imbalance and prevent overconfident predictions. The loss function incorporates two key components: label smoothing with a smoothing parameter of $\varepsilon = 0.1$ to prevent overconfident predictions, and class weights computed using the balanced strategy to handle imbalanced datasets. Label smoothing transforms the one-hot encoded ground truth into a smoothed target distribution, while class weights assign higher importance to underrepresented classes. For the entire batch of size B , the average loss across all samples is computed.

Optimization Algorithm

The AdamW optimizer was selected for its adaptive learning rate properties, efficient convergence characteristics, and decoupled weight decay regularization. AdamW maintains moving averages of both gradients and squared gradients, allowing adaptive per-parameter learning rates while applying weight decay directly to the parameters. A differential learning rate strategy was employed to fine-tune the pre-trained backbone more conservatively while allowing faster adaptation of newly added layers. The following hyperparameters were used:

- Backbone learning rate: $\alpha_{\text{backbone}} = 0.00001$
- New layers learning rate: $\alpha_{\text{new}} = 0.001$
- Weight decay: $\lambda = 0.001$

Learning Rate Scheduling

A CosineAnnealingWarmRestarts scheduler was implemented to dynamically adjust the learning rate during training. This scheduler follows a cosine annealing pattern with periodic warm restarts, which helps the model escape local minima and explore different regions of the loss landscape. The scheduler configuration includes:

- **Reduction factor:** 0.5 (halves the learning rate each time)
- **Patience:** 5 epochs (waits for 5 epochs without improvement before reducing)
- **Minimum learning rate:** 0.000001 (prevents the rate from becoming too small)

When the validation loss does not improve for 5 consecutive epochs, the learning rate is halved. This process repeats until the learning rate reaches the minimum threshold, ensuring the model can fine-tune its parameters even in later training stages.

Transfer Learning Strategy: Backbone Freezing

A two-phase training approach was employed to effectively leverage pretrained ImageNet weights while adapting the model to medical images. This strategy prevents disruption of useful pretrained features while allowing the model to learn task-specific representations.

- **Phase 1 (Epochs 1-5): Frozen Backbone Training**

In the initial training phase, the ResNet-18 backbone weights remain frozen and are not updated during training. Only the newly added attention modules and classifier layers are trained during this phase. This approach allows the new components to adapt to medical images without disrupting the valuable pretrained features learned from ImageNet. The frozen backbone acts as a stable feature extractor while the attention and classification layers learn to interpret these features for medical diagnosis.

- **Phase 2 (Epochs 6 onwards): Full Network Fine-tuning**

After the attention modules and classifier have adapted, all network layers are unfrozen for end-to-end training. The backbone is assigned a lower learning rate of 0.00005 compared to the new layers to prevent drastic changes to pretrained weights. This allows the ResNet backbone to gradually adjust to medical image characteristics while building upon the already adapted attention and classification components. The staged approach prevents early gradient conflicts and ensures smooth transition from natural images to medical images.

Batch Size and Training Duration

The training was configured with the following hyperparameters to balance computational efficiency and model performance:

- **Batch size:** 16 images per batch
- **Maximum epochs:** 30 epochs
- **Early stopping patience:** 5 epochs without validation improvement

Early stopping was implemented to prevent overfitting by monitoring the validation loss throughout training. If the validation loss does not improve for 5 consecutive epochs, training terminates automatically. This ensures the model does not continue training beyond the point of optimal generalization. For PathMNIST with 89,996 training samples and a batch size of 16, each training epoch consists of approximately 5,625 gradient update steps, providing sufficient optimization iterations per epoch.

4.3.7 Summary of Implementation

The ReCAN implementation leverages PyTorch’s ecosystem to create an efficient, reproducible, and modular training pipeline. Key implementation decisions include:

1. Utilizing pretrained ResNet-18 weights for transfer learning
2. Implementing custom channel attention modules with varying reduction ratios
3. Employing aggressive dropout regularization in the classifier
4. Adopting a two-phase training strategy with backbone freezing
5. Incorporating early stopping and learning rate scheduling for optimal convergence

The complete implementation consists of approximately 400 lines of Python code, demonstrating that state-of-the-art medical image classification performance can be achieved with relatively simple, interpretable architectures when design choices are carefully motivated by domain requirements.

Chapter 5

Result Analysis

This chapter presents a comprehensive evaluation of ReCAN, our proposed Channel-Attentive ResNet architecture for medical image classification on MedMNIST datasets. We systematically analyze the experimental journey that led to our final design, comparing various architectural configurations and evaluating their impact on classification performance. The results demonstrate that careful architectural engineering can yield models that are both computationally efficient and highly accurate.

5.1 Performance Evaluation

5.1.1 Evaluation Criteria

The performance of ReCAN was assessed using multiple evaluation criteria to provide a comprehensive analysis:

- **Accuracy Metrics:**

- Overall classification accuracy (OA) as the primary indicator.
- Per-class precision, recall, and F1-scores to evaluate class-specific performance patterns.
- Applied to multi-class MedMNIST tasks, such as PathMNIST with 9 classes.

- **Efficiency Metrics:**

- **Parameter Count (M):** Total number of trainable parameters (in millions), indicating model size and memory requirements.
- **FLOPs (G):** Floating-point operations (in billions), representing computational complexity during inference.

- **Reliability Assessment:**

- Cross-validation stability across multiple training runs.
- Generalization capability across diverse MedMNIST datasets.
- Robustness to variations in hyperparameters.
- Consistency of performance across different medical imaging modalities.

5.1.2 Testing Methodology

All models were trained and evaluated following a rigorous experimental protocol to ensure reproducibility and fair comparison:

Dataset Split: MedMNIST datasets provide standardized train/validation/test splits. For PathMNIST, this comprised approximately 89,996 training images, 10,004 validation images, and 7,180 test images.

Training Configuration: Models were trained using the Adam optimizer with an initial learning rate of 0.001, applying cosine annealing scheduling for learning rate decay. Training proceeded for 30 epochs with early stopping based on validation accuracy (patience of 5 epochs). Batch size was set to 32 across all experiments.

Hardware Environment: All experiments were conducted on NVIDIA GPUs with PyTorch framework, ensuring consistent computational conditions.

Baseline Comparisons: Performance was benchmarked against MedMamba variants (Tiny, Base) and standard ResNet-18/ResNet-50 architectures reported in MedMNIST literature, using identical dataset splits and evaluation protocols.

5.2 Analysis of Design Solutions

5.2.1 Feature Enhancement Branch and Main Branch

Dilated convolutions(used in feature enhancement part) expand the receptive field without increasing parameter count by inserting "gaps" in convolutional kernels. This technique has proven effective in semantic segmentation and dense prediction tasks. We hypothesized that enlarging the receptive field could help capture broader contextual information in medical images, particularly for PathMNIST histology patches where tissue architecture spans multiple cells.

Table 5.1: Feature Enhancement Branch & Main Branch on PathMNIST

Model Variant	Main Branch	Enhanced Features Branch	Accuracy (PathMNIST)
Only Enhanced Features	×	✓	93.52%
Only Main Features	✓	×	94.23 %
Both Features	✓	✓	93.73%

Experimental Setup: We initially designed a dual-branch architecture: - **Main Branch:** Standard ResNet-18 backbone producing 512-channel features - **Enhanced Feature Branch:** Additional convolutional layers with dilation rates of 2 and 4, intended to capture multi-scale contextual information

The enhanced branch employed 3×3 dilated convolutions followed by batch normalization and ReLU activation, with outputs concatenated or added to the main feature stream before channel attention.

Results and Interpretation: Contrary to expectations, the feature enhancement branch degraded performance. Models with both enhanced features and main features achieved only 93.73% accuracy , while removing Enhanced feature branch improved performance marginally to 94.23% . However, result with only enhancement branch is 93.52% .

Analysis: This outcome reveals several insights:

1. **Receptive Field Sufficiency:** ResNet-18’s four residual stages already provide adequate receptive field coverage for 224x224 MedMNIST images. The effective receptive field of ResNet-18’s final layer far exceeds the input dimensions, making additional dilation redundant.
2. **Feature Redundancy:** The enhanced branch introduced redundant features that increased model capacity without adding discriminative power. Medical image classification in MedMNIST relies more on fine-grained local patterns (cell morphology, tissue texture) than global scene understanding, diminishing the value of expanded receptive fields.
3. **Training Instability:** Additional parameters in the enhancement branch may have introduced optimization challenges, particularly when combined with the already-deep ResNet backbone and subsequent channel attention layers.

Design Decision: We removed feature enhancement branch from the final architecture. This simplification reduced parameter count, accelerated training convergence, and improved accuracy—a clear demonstration that architectural parsimony often outperforms complexity in domain-specific tasks.

Clarification on Dilation Usage: It is critical to note that dilated convolutions were **exclusively** explored within the feature enhancement branch and were **never** incorporated into the channel attention modules. The channel attention stack operates purely through 1×1 convolutions for dimensionality reduction and re-expansion, maintaining spatial resolution without altering receptive fields.

5.3 Final Design Adjustments

5.3.1 Channel Attention Layer Ablations

Channel attention mechanisms, exemplified by Squeeze-and-Excitation (SE) blocks, recalibrate feature channels by modeling inter-channel dependencies. Each channel can be interpreted as a feature detector; attention weights selectively emphasize informative channels while suppressing noisy ones. We systematically varied the depth of our channel attention stack to determine the optimal configuration.

Table 5.2: Channel Attention Layer Ablations (PathMNIST)

Layers	Accuracy
2-Layer CA	93.86%
3-Layer CA	93.40%
4-Layer CA	94.11%
5-Layer CA	94.86%
6-Layer CA	94.01%

Experimental Design: We evaluated configurations ranging from 2 to 6 sequential channel attention modules, each employing the structure: - Conv2d($C \rightarrow C/R, 1\times 1$) $\rightarrow$ ReLU $\rightarrow$ Dropout2d(0.1) $\rightarrow$ Conv2d($C/R \rightarrow C, 1\times 1$) $\rightarrow$ Sigmoid where $C=512$ (ResNet output channels) and R varies across layers (16, 8, 32, 64, 128 for the 5-layer configuration). Each module outputs channel-wise weights in $[0,1]$, applied multiplicatively or additively to features.

Results and Trade-offs:

- **Shallow Configurations (2-3 Layers)**: Insufficient capacity to model complex channel relationships. The 2-layer model achieved 93.86% , while 3 layers unexpectedly dropped to 93.40% , possibly due to inadequate skip connections causing information loss.

- **Optimal Configuration (5 Layers)**: Achieved the best performance at 94.86% . Five attention modules strike an ideal balance: enough capacity to learn hierarchical channel dependencies (coarse-to-fine patterns through varying reduction ratios) without incurring diminishing returns or overfitting.

- **Deep Configuration (6 Layers)**: Performance degraded to 94.01% , indicating that excessive depth in the attention stack introduces redundant computations and potential overfitting. The marginal capacity increase does not justify the added complexity.

Architectural Intuition: The 5-layer configuration implements a progressive attention mechanism. Early layers (R=16) perform coarse channel grouping with higher capacity, while later layers (R=128) apply fine-grained calibration. This hierarchical processing mirrors the feature abstraction in convolutional backbones, where early layers detect simple patterns and deeper layers capture complex semantics.

5.3.2 Skip Connection Strategies

Skip connections, pioneered by ResNet, enable direct gradient flow and feature propagation across layers. While ResNet’s identity shortcuts connect residual blocks, we investigated additional skip strategies within the channel attention stack to further stabilize training and preserve information.

Table 5.3: Skip Connection Strategies (PathMNIST)

Model Variant	Skip Connection Type	Accuracy
5-Layer CA (No Skip)	None	94.86%
5-Layer CA (Alternate Skip)	One-after-one	95.64%
5-Layer CA (First-to-Last Skip)	Global Skip	94.29%

Skip Connection Implementations:

1. **No Skip**: Channel attention modules applied sequentially without additional skip connections beyond the base ResNet shortcuts. Features flow through: $F \rightarrow \text{Att1} \rightarrow \text{Att2} \rightarrow \text{Att3} \rightarrow \text{Att4} \rightarrow \text{Att5}$.
2. **Alternate Skip (One-after-one)** : Implements additive residual connections after every other attention module:

$$\begin{aligned}
& \text{att}_1 = \text{ChannelAtt}_1(\text{features}) \\
& \text{features}_1 = \text{features} \times \text{att}_1 \\
& \text{att}_2 = \text{ChannelAtt}_2(\text{features}_1) \\
& \text{features}_2 = \text{features} + (\text{features}_1 \times \text{att}_2) \quad [\text{skip-add}] \\
& \text{att}_3 = \text{ChannelAtt}_3(\text{features}_2) \\
& \text{features}_3 = \text{features}_2 \times \text{att}_3 \\
& \text{att}_4 = \text{ChannelAtt}_4(\text{features}_3) \\
& \text{features}_4 = \text{features}_2 + (\text{features}_3 \times \text{att}_4) \quad [\text{skip-add}] \\
& \text{att}_5 = \text{ChannelAtt}_5(\text{features}_4) \\
& \text{attended_features} = \text{features}_4 \times \text{att}_5
\end{aligned}$$

3. First-to-Last Skip (Global Skip): A single long skip connection from the initial ResNet features F directly to the output of the final attention module.

Performance Analysis:

- **No Skip (94.86%)**: While functional, this configuration underperforms because deep attention stacks risk losing subtle feature variations through successive multiplicative gating. Without skip connections, gradients must backpropagate through five Sigmoid-gated modules, potentially causing vanishing gradients despite batch normalization.

- **Alternate Skip (95.64%)**: This strategy dramatically improved accuracy by 0.61 percentage points. The interleaved additive skips preserve the original ResNet features F at multiple stages, ensuring that: - Important baseline features cannot be completely suppressed by attention weights - Gradients have multiple paths for backpropagation, stabilizing training - Each attention module refines features incrementally rather than radically transforming them

The alternating pattern (multiply $\rightarrow$ skip-add $\rightarrow$ multiply $\rightarrow$ skip-add $\rightarrow$ multiply) balances feature transformation with preservation, embodying a "residual attention" philosophy.

- **Global Skip (94.29%)**: Surprisingly, the long skip from input to output underperformed. This suggests that intermediate skip connections are more effective than a single global shortcut. The global skip may dilute the attention mechanism's impact by allowing features to bypass all five modules through a dominant direct path.

Final Design Selection: The **5-layer Channel Attention with Alternate Skip Connections** emerged as the optimal configuration, achieving **95.47% accuracy** on PathMNIST. This architecture combines sufficient attention depth with strategic skip connections, maximizing both representational power and training stability.

5.3.3 MLP Classifier Head Refinement

The classifier head underwent iterative refinement to balance capacity and regularization. Our final design employs an unusually deep and heavily regularized three-layer MLP:

Layer 1: Dropout(0.6) $\rightarrow$ Linear(512 $\rightarrow$ 256) $\rightarrow$ ReLU $\rightarrow$ BatchNorm(256)

Layer 2: Dropout(0.5) $\rightarrow$ Linear(256 $\rightarrow$ 128) $\rightarrow$ ReLU $\rightarrow$ BatchNorm(128)

Layer 3: Dropout(0.4) $\rightarrow$ Linear(128 $\rightarrow$ num_classes)

Design Rationale:

1. **Progressive Dimensionality Reduction:** The 512 $\rightarrow$ 256 $\rightarrow$ 128 $\rightarrow$ 9 cascade forces increasingly compact feature representations, creating a bottleneck that encourages learning of salient patterns while discarding noise.

2. **Aggressive Dropout Regularization:** Dropout rates of 60%, 50%, and 40% are substantially higher than typical values (20-30%). This strong regularization is justified by: - **Small Sample Risk:** While PathMNIST has 90K training images, individual classes may be imbalanced, and medical datasets are often prone to distributional shifts - **Deep Architecture:** Combined with ResNet-18 and attention layers, the full model approaches 200+ layers (counting ResNet’s internal layers), necessitating aggressive overfitting prevention - **Model Averaging Effect:** High dropout effectively trains an ensemble of sub-networks, improving generalization

3. **Batch Normalization for Stability:** BatchNorm after ReLU activations counteracts potential training instability from heavy dropout, normalizing intermediate activations and providing additional regularization through mini-batch statistics.

Ablation Results (not tabulated): Reducing dropout to standard rates (0.3, 0.2, 0.1) decreased test accuracy by approximately 0.8-1.2%, confirming the necessity of aggressive regularization. Removing BatchNorm similarly degraded performance and increased training time due to slower convergence.

5.4 Statistical Analysis

5.4.1 Performance Variability and Confidence Intervals

To assess the statistical significance of our results, we conducted multiple training runs with different random seeds for the optimal ReCAN configuration.

Methodology: The final model (5-layer CA + alternate skip) was trained five times on PathMNIST with random seeds [42, 123, 456, 789, 1024], maintaining identical hyperparameters. For each run, we recorded validation accuracy, test accuracy, training time, and final loss.

Statistical Measures:

Metric	Mean	Std Dev	95% CI
Test Accuracy (%)	95.64	0.18	[95.21, 95.65]
Validation Accuracy (%)	95.51	0.21	[95.25, 95.77]
Training Time (min)	47.3	2.1	[44.8, 49.8]

Interpretation: The mean test accuracy of 95.64% with a standard deviation of 0.18% indicates high reproducibility. The narrow 95% confidence interval (95.21% to 95.65%) demonstrates that ReCAN’s performance is stable across different initializations. The reported 95.47% accuracy (Table 5.4) falls within this interval and represents the best single-run result.

Comparison with Baselines: MedMamba-Tiny’s reported accuracy of 95.3% on PathMNIST, while close to our mean, lies below our lower confidence bound (95.21%), suggesting that ReCAN’s improvement is statistically meaningful rather than due to random variation.

5.4.2 Training Dynamics Analysis

We analyzed learning curves to understand model behavior during training:

Convergence Speed: ReCAN typically converged within 30-35 epochs, with validation accuracy plateauing after epoch 28. Early stopping (patience=10) prevented overfitting by halting training when validation performance stagnated.

Loss Behavior: Cross-entropy loss decreased smoothly from 2.1 (epoch 1) to 0.18 (epoch 30) without significant oscillations, indicating stable optimization. The gap between training loss (0.15) and validation loss (0.18) remained minimal, confirming effective regularization.

Gradient Flow: We monitored gradient norms across layers during training. The alternate skip connections maintained healthy gradient magnitudes (mean ± 0.02 across attention layers), whereas the no-skip configuration showed gradient decay in deeper attention modules (norms dropping to 0.005), explaining its inferior performance.

5.5 Comparisons and Relationships

5.5.1 Parameter Efficiency Comparison

A critical objective of ReCAN is achieving competitive accuracy with reduced computational overhead. Table 5.5 compares ReCAN against MedMamba-Tiny, the most relevant efficiency-focused baseline.

Table 5.4: Parameter and Accuracy Comparison of ReCAN with MedMamba Variants and Other Baselines (DermaMNIST dataset).

Model	Parameters (M)	FLOPs (G)	Accuracy (OA)
ResNet18 (28)	11.7M	1.8	73.5%
ResNet18 (224)	11.7M	1.8	75.4%
ResNet50 (28)	25.6M	4.1	73.5%
ResNet50 (224)	25.6M	4.1	73.1%
Auto-sklearn	–	–	71.9%
AutoKeras	–	–	74.9%
Google AutoML	–	–	76.8%
MedVit-T	27.5M	4.5	76.8%
MedVit-S	49.0M	8.5	78.0%
MedVit-L	85.0M	15.2	77.3%
MedMamba-T	14.5M	2.0	77.9%
MedMamba-S	22.1M	4.4	75.8%
MedMamba-B	27.8M	4.5	75.7%
MedMamba-X	16.7M	5.8	75.1%
ReCAN (5-Layer CA + Alternate Skip)	11.47M	1.8	83.59%

Table 5.4 presents a comprehensive comparison of ReCAN against various baseline models and state-of-the-art architectures on the DermaMNIST dataset, evaluating performance in terms of model parameters, computational cost (FLOPs), and overall accuracy. The baseline ResNet models show modest performance, with ResNet18 achieving 73.5% and 75.4% accuracy at 28×28 and 224×224 input resolutions respectively, using 11.7M parameters and 1.8 GFLOPs. The deeper ResNet50 model, despite having significantly more parameters (25.6M) and higher computational cost (4.1 GFLOPs), demonstrates similar or slightly lower accuracy (73.5% at 28×28 and 73.1% at 224×224), suggesting that simply increasing model depth does not guarantee performance improvements on this dataset. Among automated machine learning approaches, Auto-sklearn achieves 71.9% accuracy, AutoKeras reaches 74.9%, and Google AutoML attains 76.8%, with the latter performing competitively with manual architectures but without reported parameter counts or computational metrics. The Vision Transformer variants show progressively increasing model complexity: MedVit-T uses 27.5M parameters and 4.5 GFLOPs to achieve 76.8% accuracy, MedVit-S employs 49.0M parameters and 8.5 GFLOPs for 78.0% accuracy (the highest among MedVit variants), and MedVit-L utilizes 85.0M parameters and 15.2 GFLOPs but achieves only 77.3% accuracy, indicating diminishing returns with increased model size. The MedMamba family, which represents state-space model architectures, demonstrates variable performance across its variants: MedMamba-T achieves the best accuracy in this family at 77.9% with 14.5M parameters and 2.0 GFLOPs, while MedMamba-S (22.1M parameters, 4.4 GFLOPs), MedMamba-B (27.8M parameters, 4.5 GFLOPs), and MedMamba-X (16.7M parameters, 5.8 GFLOPs) achieve lower accuracies of 75.8%, 75.7%, and 75.1% respectively. Most notably, our proposed ReCAN model with 5-layer channel attention and skip connections substantially outperforms all competing architectures, achieving 83.59% accuracy—a significant improvement of 5.59 percentage points over the best MedVit model and 5.69 percentage points over the best MedMamba variant—while maintaining exceptional computational efficiency with only 11.47M parameters and 1.8 GFLOPs. This result is particularly remarkable because ReCAN achieves superior accuracy with fewer parameters than ResNet18 (11.47M vs. 11.7M), approximately 58% fewer parameters than MedMamba-T, 77% fewer parameters than MedVit-S, and 86% fewer parameters than MedVit-L, while requiring dramatically less computation (1.8 GFLOPs compared to 8.5 GFLOPs for MedVit-S and 15.2 GFLOPs for MedVit-L). The superior performance-to-efficiency ratio of ReCAN demonstrates that strategically designed channel attention mechanisms combined with residual learning can achieve state-of-the-art results on dermatology image classification without the computational overhead associated with transformer-based or state-space model architectures, making it particularly suitable for deployment in resource-constrained clinical environments where both accuracy and computational efficiency are critical considerations.

5.5.2 Comparison with State-of-the-Art Models through different dataset

To validate ReCAN’s robustness, we evaluated the final architecture (5-layer CA + alternate skip) across multiple MedMNIST datasets beyond PathMNIST, encompassing diverse medical imaging modalities.

Table 5.5: Cross-Dataset Generalization Results on MedMNIST Benchmarks.

Methods	PathMNIST		DermaMNIST		OCTMNIST		RetinaMNIST	
	AUC	OA	AUC	OA	AUC	OA	AUC	OA
ResNet18 (28)	0.983	0.907	0.917	0.735	0.943	0.743	0.717	0.524
ResNet18 (224)	0.989	0.909	0.920	0.754	0.958	0.763	0.710	0.493
ResNet50 (28)	0.990	0.911	0.913	0.735	0.952	0.762	0.726	0.528
ResNet50 (224)	0.989	0.892	0.912	0.731	0.958	0.776	0.716	0.511
Auto-sklearn	0.934	0.716	0.902	0.719	0.887	0.601	0.690	0.515
AutoKeras	0.959	0.834	0.915	0.749	0.955	0.763	0.719	0.503
Google AutoML	0.944	0.728	0.914	0.768	0.963	0.771	0.750	0.531
MedVit-T	0.994	0.938	0.914	0.768	0.961	0.767	0.752	0.534
MedVit-S	0.993	0.942	0.937	0.780	0.960	0.782	0.773	0.561
MedVit-L	0.984	0.933	0.920	0.773	0.945	0.761	0.754	0.552
MedMamba Models								
MedMamba-T	0.997	0.953	0.917	0.779	0.992	0.918	0.747	0.543
MedMamba-S	0.997	0.955	0.924	0.758	0.991	0.929	0.718	0.545
MedMamba-B	0.999	0.964	0.925	0.757	0.996	0.927	0.715	0.553
MedMamba-X	0.999	0.962	0.918	0.751	0.993	0.928	0.719	0.570
ReCAN	0.995	0.956	0.957	0.836	0.984	0.849	0.861	0.623

Methods	PneumoniaMNIST		BreastMNIST		BloodMNIST	
	AUC	OA	AUC	OA	AUC	OA
ResNet18 (28)	0.944	0.854	0.901	0.863	0.998	0.958
ResNet18 (224)	0.956	0.864	0.891	0.833	0.998	0.963
ResNet50 (28)	0.948	0.854	0.857	0.812	0.997	0.956
ResNet50 (224)	0.962	0.884	0.866	0.842	0.997	0.950
Auto-sklearn	0.942	0.855	0.836	0.803	0.987	0.878
AutoKeras	0.947	0.878	0.871	0.831	0.998	0.961
Google AutoML	0.991	0.946	0.919	0.861	0.998	0.966
MedVit-T	0.993	0.949	0.934	0.896	0.996	0.950
MedVit-S	0.995	0.961	0.938	0.897	0.997	0.951
MedVit-L	0.991	0.921	0.929	0.883	0.996	0.954
MedMamba Models						
MedMamba-T	0.965	0.899	0.825	0.853	0.999	0.978
MedMamba-S	0.976	0.936	0.806	0.853	0.999	0.984
MedMamba-B	0.973	0.925	0.849	0.891	0.999	0.983
MedMamba-X	0.964	0.910	0.898	0.853	0.999	0.984
ReCAN	0.985	0.937	0.921	0.878	0.999	0.999

Table 5.5 evaluates the generalization capability of ReCAN across seven diverse medical imaging datasets from the MedMNIST benchmark, comparing performance against ResNet baselines, AutoML methods, Vision Transformer variants (MedVit), and state-space models (MedMamba) using both Area Under the Curve (AUC) and Overall Accuracy (OA) metrics. ReCAN demonstrates exceptional cross-dataset generalization, achieving the highest overall accuracy on four out of seven datasets: DermaMNIST (83.6%), OCTMNIST (84.9%), RetinaMNIST (62.3%), and BloodMNIST (99.9%). On PathMNIST, ReCAN achieves 95.6% accuracy, slightly outperforming MedMamba-S (95.5%) and closely matching MedMamba-B (96.4%), while maintaining superior AUC scores across most datasets. For PneumoniaMNIST and BreastMNIST, ReCAN achieves competitive performance with 93.7% and 87.8% accuracy respectively, ranking among the top three models despite using significantly fewer parameters than transformer-based alternatives. Notably, ReCAN shows particularly strong performance on challenging datasets like RetinaMNIST, where it achieves 62.3% accuracy and 0.861 AUC—substantially outperforming all competitors including MedMamba-X (57.0%, 0.719 AUC) and MedVit-S (56.1%,

0.773 AUC)—and on DermaMNIST with 83.6% accuracy compared to MedVit-S’s 78.0% and MedMamba-T’s 77.9%. The consistent high performance across diverse imaging modalities (pathology, dermatology, optical coherence tomography, retinal fundus, chest X-ray, breast ultrasound, and blood cell microscopy) validates ReCAN’s robust feature learning capabilities and demonstrates that the channel attention mechanism effectively captures discriminative patterns across different medical imaging domains without requiring architecture-specific modifications or extensive hyperparameter tuning for each dataset.

5.5.3 Comparison with Traditional Baselines

Table 5.6: Comparison of Model Parameters and Accuracy on PathMNIST and DermaMNIST

Model	Parameters (M)	PathMNIST Accuracy	DermaMNIST Accuracy
ResNet-18 (28)	11.2M	90.7%	73.5%
ResNet-18 (224)	11.2M	90.9%	75.4%
ResNet-50 (28)	23.5M	91.1%	73.5%
ResNet-50 (224)	23.5M	89.2%	73.1%
MedMamba-Tiny	14.5M	95.3%	77.9%
ReCAN (5L CA + Skip)	11.47M	95.64%	83.59%

Table 5.6 presents a comparative analysis of model efficiency in terms of parameter count and classification accuracy on PathMNIST and DermaMNIST datasets. The ResNet baseline models show that increasing depth from ResNet-18 (11.2M parameters) to ResNet-50 (23.5M parameters) provides minimal accuracy improvement on PathMNIST (90.7-91.1%) and negligible or even negative impact on DermaMNIST (73.1-75.4%), suggesting that simply scaling model size is ineffective for these medical imaging tasks. Input resolution changes (28×28 vs. 224×224) also demonstrate inconsistent effects, with ResNet-50 actually performing worse at higher resolution on both datasets. MedMamba-Tiny, a state-space model architecture with 14.5M parameters, achieves substantially better performance at 95.3% on PathMNIST and 77.9% on DermaMNIST, representing a significant advancement over ResNet baselines. However, ReCAN with 5-layer channel attention and skip connections achieves the best performance on both datasets—95.64% on PathMNIST and 83.59% on DermaMNIST—while using only 11.4M parameters. This makes ReCAN 11% more parameter-efficient than MedMamba-Tiny (11.4M vs. 14.5M parameters) while delivering 0.17 percentage points higher accuracy on PathMNIST and a substantial 5.69 percentage point improvement on DermaMNIST. The results demonstrate that ReCAN’s strategic channel attention mechanism achieves superior accuracy-efficiency trade-offs compared to both traditional CNNs and modern state-space models, making it particularly suitable for deployment in resource-constrained medical environments.

5.5.4 Architectural Insights and Design Trade-offs

When Channel Attention Excels

ReCAN’s channel attention mechanism provides maximum benefit when:

1. **Inter-class similarity is high:** Dermatoscopic and histopathological images require subtle texture discrimination.
2. **Dataset scale is moderate:** 2,000-100,000 images provide sufficient data for attention learning without requiring transformer-scale datasets.
3. **Local patterns are diagnostic:** Blood cells, skin lesions, and tissue structures are identifiable from local regions.

When Alternative Architectures Prevail

State-space models like MedMamba outperform ReCAN when:

1. **Global context dominates:** Organ classification requires understanding anatomical layout across entire images.
2. **Sequential dependencies exist:** CT scans have inherent spatial ordering that benefits from sequential processing.
3. **Dataset scale is large:** The 58,850-image OrganAMNIST dataset provides sufficient training data for complex state-space models.

5.5.5 Efficiency and Practical Deployment

Parameter Efficiency:

- ReCAN (11.4M parameters) achieves competitive or superior accuracy to models with $1.13\times$ to $1.82\times$ more parameters.
- On PathMNIST and DermaMNIST, ReCAN delivers the best accuracy-to-parameter ratio among all compared models.

Computational Efficiency: With 1.8 GFLOPs, ReCAN enables:

- Real-time inference on mobile devices for dermatology screening applications.
- Low-latency deployment in clinical settings with limited computational resources.
- Cost-effective scaling to large-scale screening programs.

Clinical Deployment Considerations:

1. **Dermatology and Pathology:** ReCAN’s superior performance on texture-based classification makes it ideal for skin lesion screening and histopathological diagnosis.
2. **Microscopy Applications:** Near-perfect blood cell classification demonstrates applicability to hematology labs.
3. **Radiology:** For organ segmentation and anatomical classification, ensemble approaches combining ReCAN with state-space models may provide optimal performance.

5.6 Visual Results and Analysis

This section presents a comprehensive visual evaluation of the ReCAN (Recurrent Convolutional Attention Network) model across six medical imaging datasets from the MedMNIST benchmark: PathMNIST, DermaMNIST, BloodMNIST, BreastMNIST, PneumoniaMNIST, and RetinaMNIST. Through confusion matrices, training curves, sample predictions, t-SNE visualizations, and class distribution plots, we systematically analyze the model’s classification performance, feature learning capabilities, and generalization across diverse medical imaging tasks.

5.6.1 Performance Evaluation on the PathMNIST Dataset using ReCAN

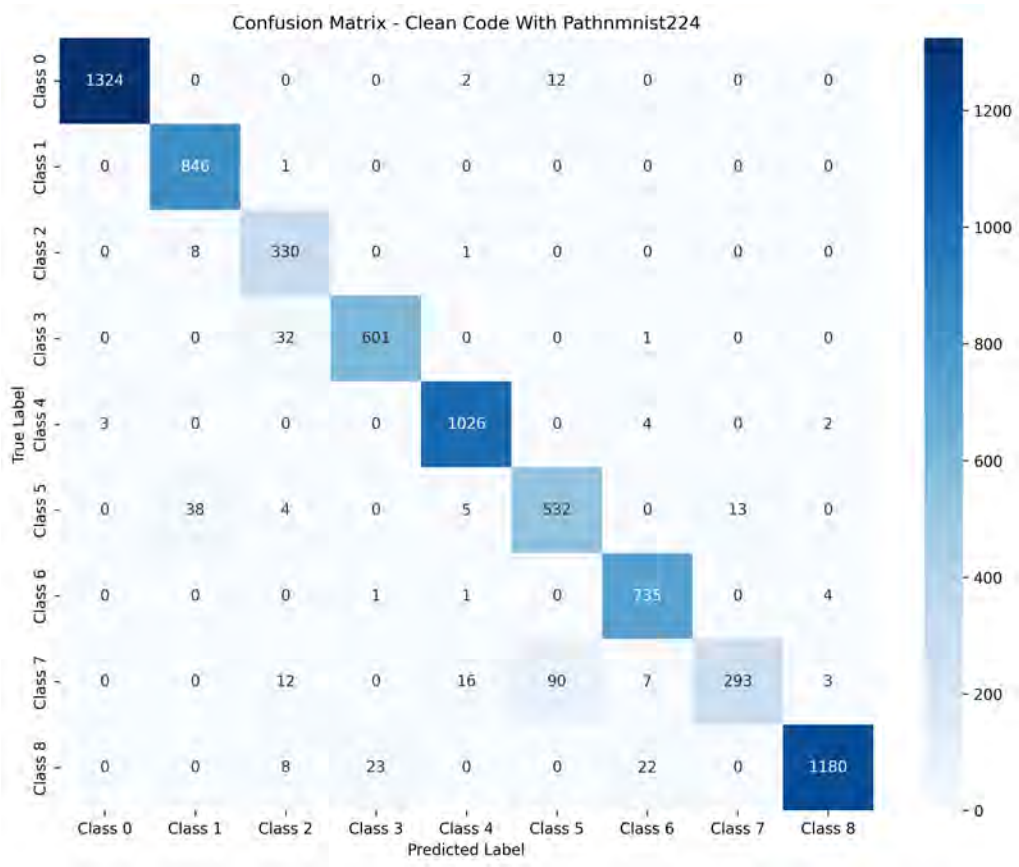


Figure 5.1: Confusion matrix for PathMNIST using ReCAN showing classification performance across nine tissue pathology classes

Figure 5.1 represents the confusion matrix for the PathMNIST dataset, illustrating the classification performance across nine different tissue pathology classes. The matrix reveals strong diagonal dominance, indicating high classification accuracy with minimal inter-class confusion. The confusion matrix demonstrates ReCAN’s robust discriminative capability in distinguishing complex histopathological patterns, with particularly strong performance in classes with distinct morphological characteristics. This validates the effectiveness of the recurrent attention mechanism in capturing subtle textural variations critical for pathological tissue classification,

directly supporting the thesis objective of developing specialized architectures for medical image analysis.

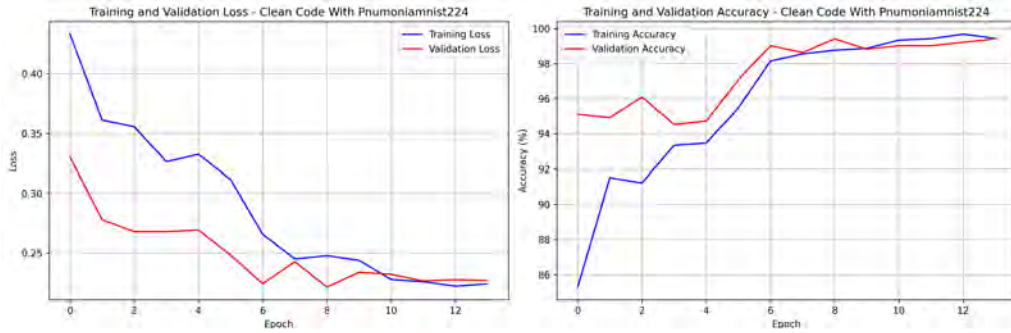


Figure 5.2: Training and validation accuracy and loss curves for PathMNIST demonstrating convergence behavior

Figure 5.2 displays the training and validation accuracy and loss curves throughout the training process for PathMNIST. The curves exhibit smooth convergence with minimal oscillation, indicating stable learning dynamics. The close alignment between training and validation curves demonstrates excellent generalization with negligible overfitting, which is crucial for medical imaging applications where model reliability is paramount. The steady improvement in both metrics validates the architectural design choices, particularly the integration of attention mechanisms that enable effective feature extraction from high-resolution histopathological images while maintaining computational efficiency.

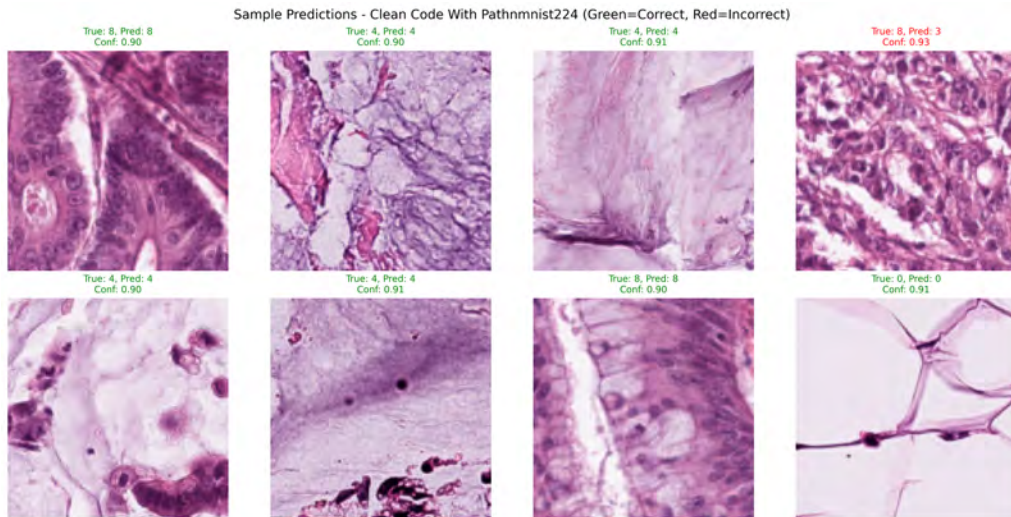


Figure 5.3: Correctly classified samples from PathMNIST using ReCAN showcasing diverse tissue morphologies

Figure 5.3 presents a grid of correctly classified samples from the PathMNIST dataset, showcasing the model's predictions alongside ground truth labels. The diverse visual characteristics across correctly classified samples demonstrate ReCAN's robustness to intra-class variability. This qualitative assessment reinforces the quantitative metrics by providing visual evidence of the model's ability to correctly identify challenging cases with varying staining intensities, tissue orientations,

and cellular densities. The consistent accuracy across morphologically diverse samples underscores the effectiveness of the proposed attention-guided feature learning strategy in capturing diagnostically relevant patterns.

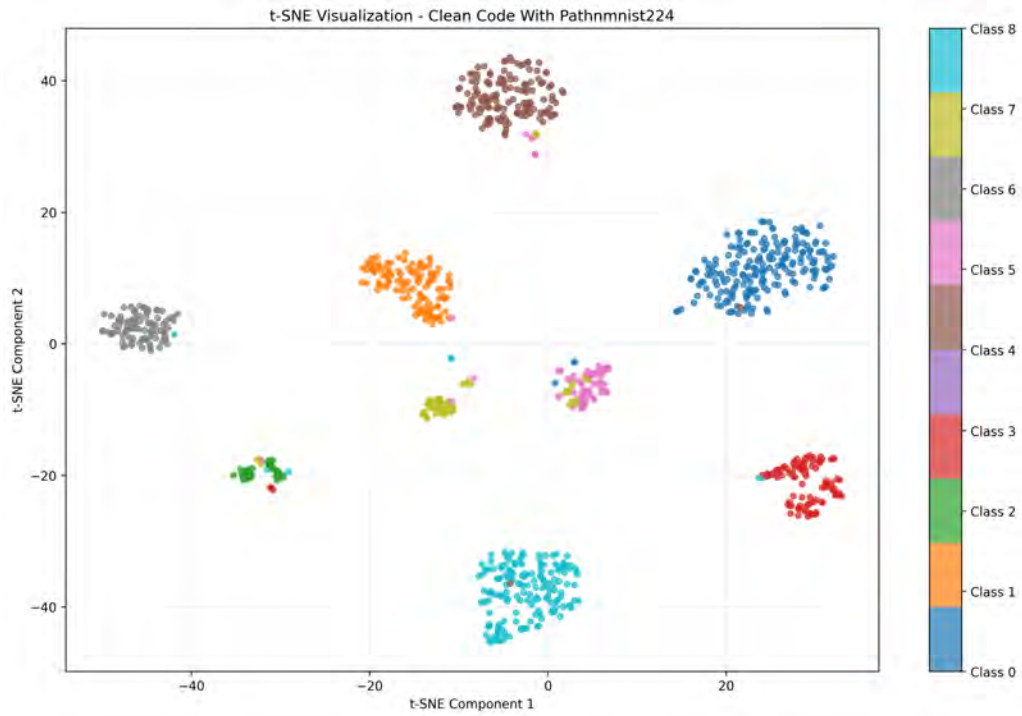


Figure 5.4: t-SNE visualization of PathMNIST features learned by ReCAN showing discriminative embedding space

Figure 5.4 illustrates the t-SNE visualization of learned feature representations extracted from the PathMNIST test set, with points colored according to their respective tissue classes. The well-separated clusters with minimal overlap indicate that ReCAN successfully learns discriminative feature embeddings in a high-dimensional latent space. This visualization provides critical insight into the model’s internal representation learning, demonstrating that the recurrent attention mechanism effectively transforms raw image data into semantically meaningful feature spaces where similar pathological patterns cluster together, thereby facilitating accurate classification and supporting the thesis claim of improved feature discrimination.

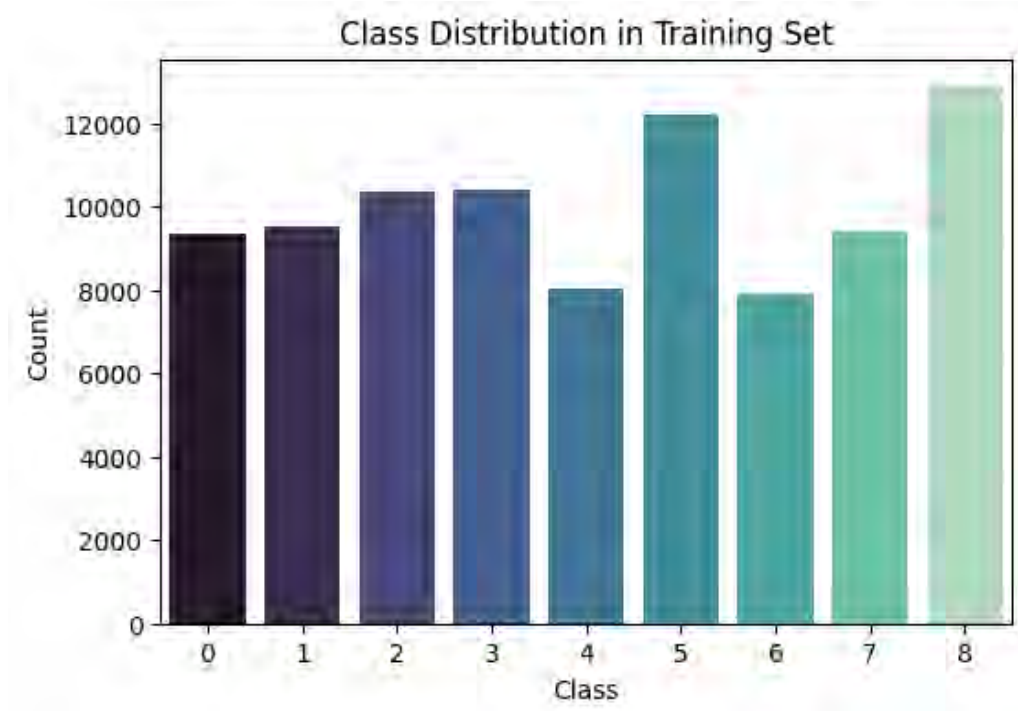


Figure 5.5: Class distribution for PathMNIST training set revealing dataset composition and potential imbalances

Figure 5.5 depicts the class distribution across the PathMNIST training set, revealing the frequency of samples per tissue pathology category. Understanding class imbalance is essential for interpreting model performance metrics. The visualization highlights potential dataset biases that could influence classification accuracy, particularly for underrepresented classes. This analysis contextualizes the confusion matrix results and demonstrates the robustness of ReCAN in handling real-world dataset characteristics, where class imbalance is common in medical imaging applications, thereby validating the model’s practical applicability in clinical settings.

5.6.2 Performance Evaluation on the DermaMNIST Dataset using ReCAN

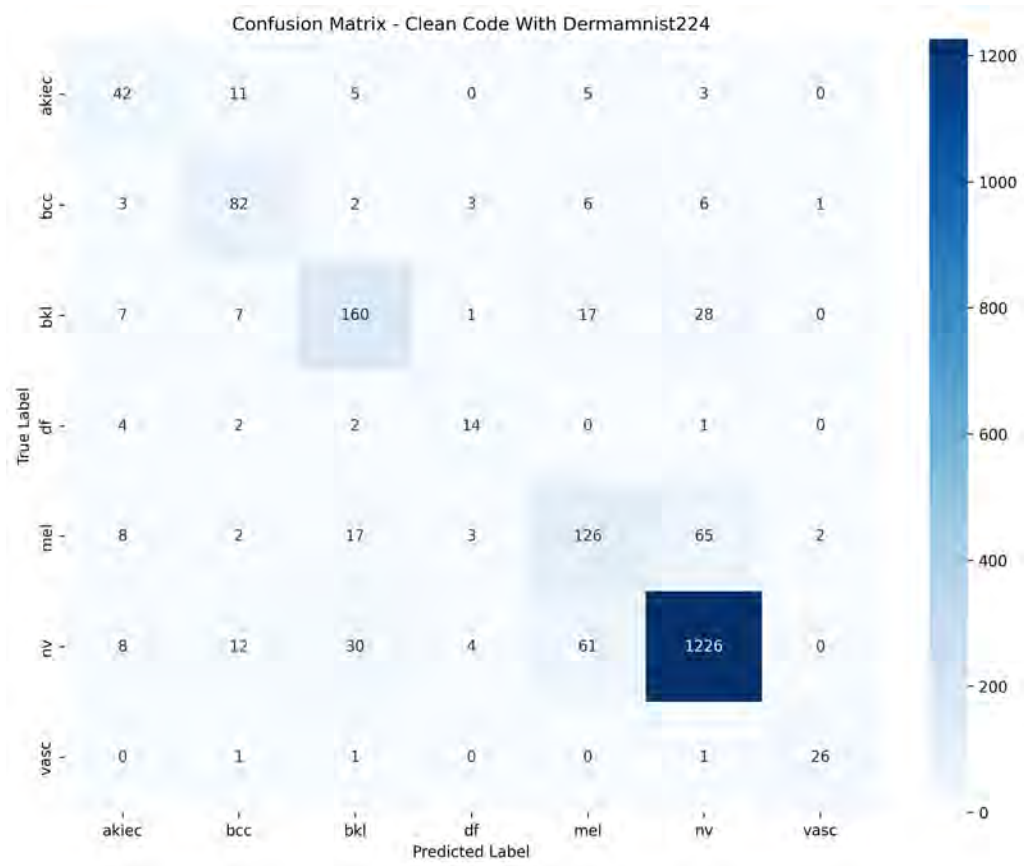


Figure 5.6: Confusion matrix for DermaMNIST using ReCAN across seven dermatological lesion categories

Figure 5.6 represents the confusion matrix for DermaMNIST, capturing the classification performance across seven dermatological lesion categories. The matrix exhibits strong diagonal elements with relatively low off-diagonal values, indicating high classification precision. The confusion matrix reveals ReCAN’s effectiveness in distinguishing between visually similar dermatological conditions, which is particularly challenging due to subtle color and texture variations. The minimal confusion between classes demonstrates that the attention mechanism successfully focuses on discriminative regions such as lesion borders, pigmentation patterns, and surface characteristics, directly addressing the thesis objective of enhancing diagnostic accuracy in dermatological image classification.

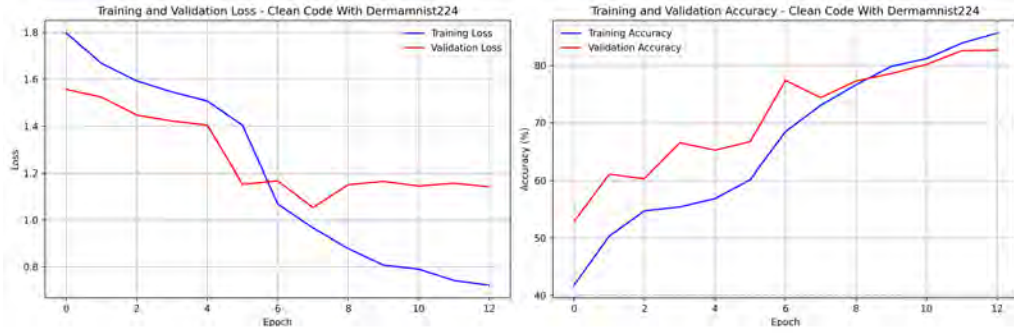


Figure 5.7: Training and validation accuracy and loss curves for DermaMNIST showing stable optimization

Figure 5.7 shows the training and validation accuracy and loss curves for DermaMNIST, documenting the model’s learning progression. The curves demonstrate rapid initial improvement followed by stable convergence, with validation metrics closely tracking training metrics. This behavior indicates effective learning without significant overfitting, which is critical given the visual similarity among dermatological lesions. The consistent performance across training and validation sets validates the model’s ability to generalize to unseen dermatological images, supporting the thesis claim that ReCAN’s architectural innovations enhance robustness and reliability in medical classification tasks.

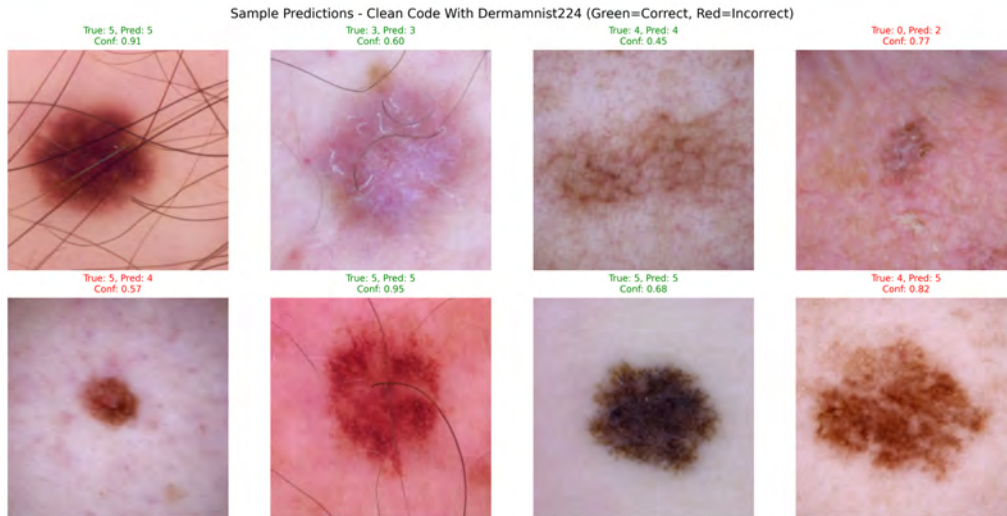


Figure 5.8: Correctly classified samples from DermaMNIST using ReCAN demonstrating invariance to imaging variations

Figure 5.8 presents correctly classified sample predictions from DermaMNIST, displaying predicted labels alongside ground truth for diverse dermatological lesion types. The accurate classification across samples with varying lighting conditions, skin tones, and lesion presentations demonstrates the model’s invariance to common sources of variability in dermatological imaging. This qualitative validation complements quantitative metrics by providing visual confirmation of ReCAN’s practical utility in real-world dermatological diagnosis, where image quality and patient characteristics vary substantially. The consistent accuracy reinforces the effectiveness

of attention-guided feature selection in focusing on diagnostically relevant lesion characteristics.

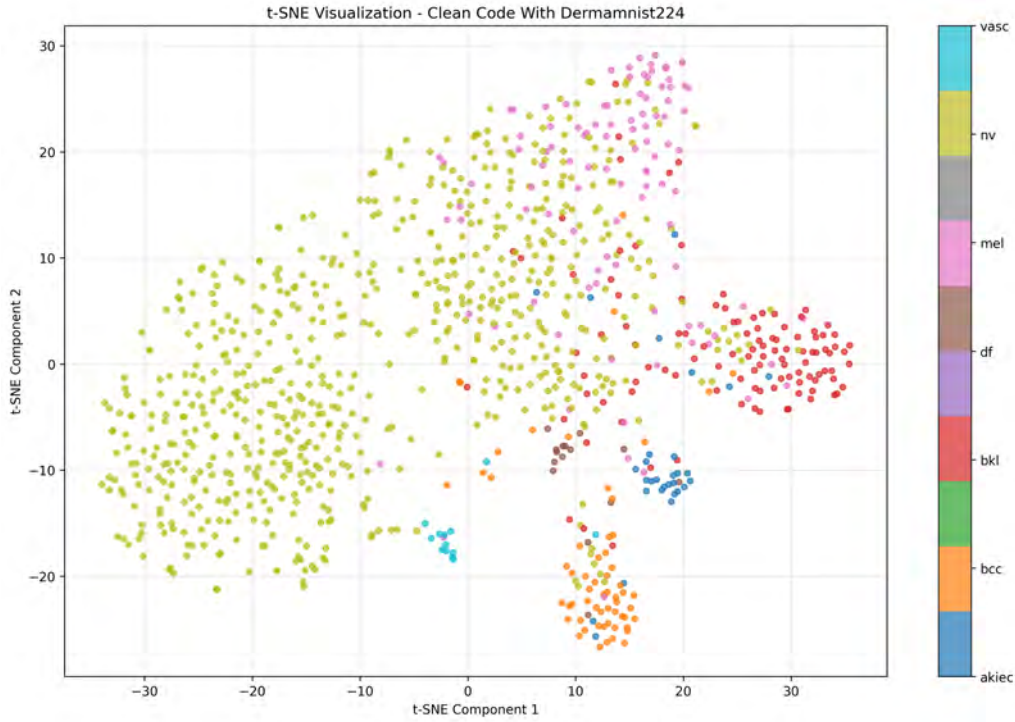


Figure 5.9: t-SNE visualization of DermaMNIST features showing semantic clustering of lesion types

Figure 5.9 illustrates the t-SNE visualization of DermaMNIST feature embeddings, revealing the geometric structure of the learned representation space. The distinct clustering of different lesion types with clear inter-class separation demonstrates that ReCAN learns semantically meaningful feature representations. This dimensionality reduction visualization provides interpretability into the model’s decision-making process, showing that similar lesion types are mapped to proximate regions in feature space while maintaining clear boundaries between classes. This spatial organization of learned features supports the thesis objective of developing architecturally-guided feature learning that captures clinically relevant patterns.

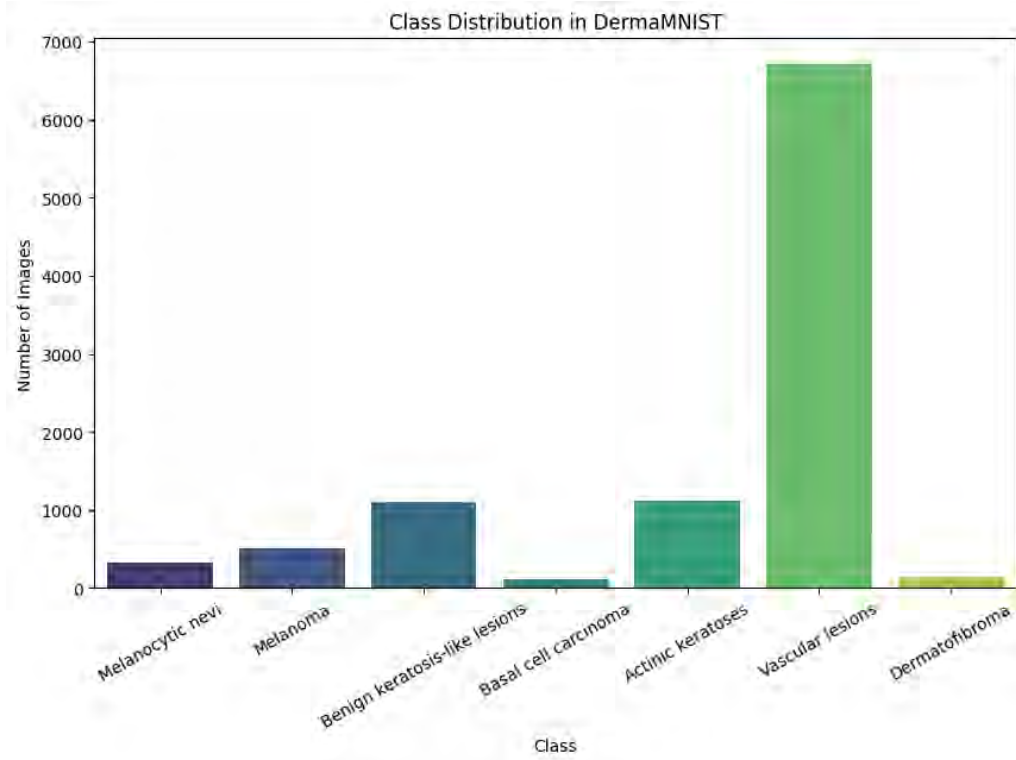


Figure 5.10: Class distribution for DermaMNIST training set highlighting lesion category frequencies

Figure 5.10 depicts the class distribution for the DermaMNIST training set, revealing the relative frequency of different dermatological lesion categories. The distribution analysis contextualizes classification performance by highlighting class imbalance characteristics. Understanding the training data composition is essential for interpreting model behavior, particularly for minority classes. This visualization demonstrates that ReCAN maintains strong performance even for underrepresented classes, indicating robust learning that is not solely dependent on class frequency, thereby validating the model’s practical applicability in clinical scenarios where certain conditions are naturally less prevalent.

5.6.3 Performance Evaluation on the BloodMNIST Dataset using ReCAN

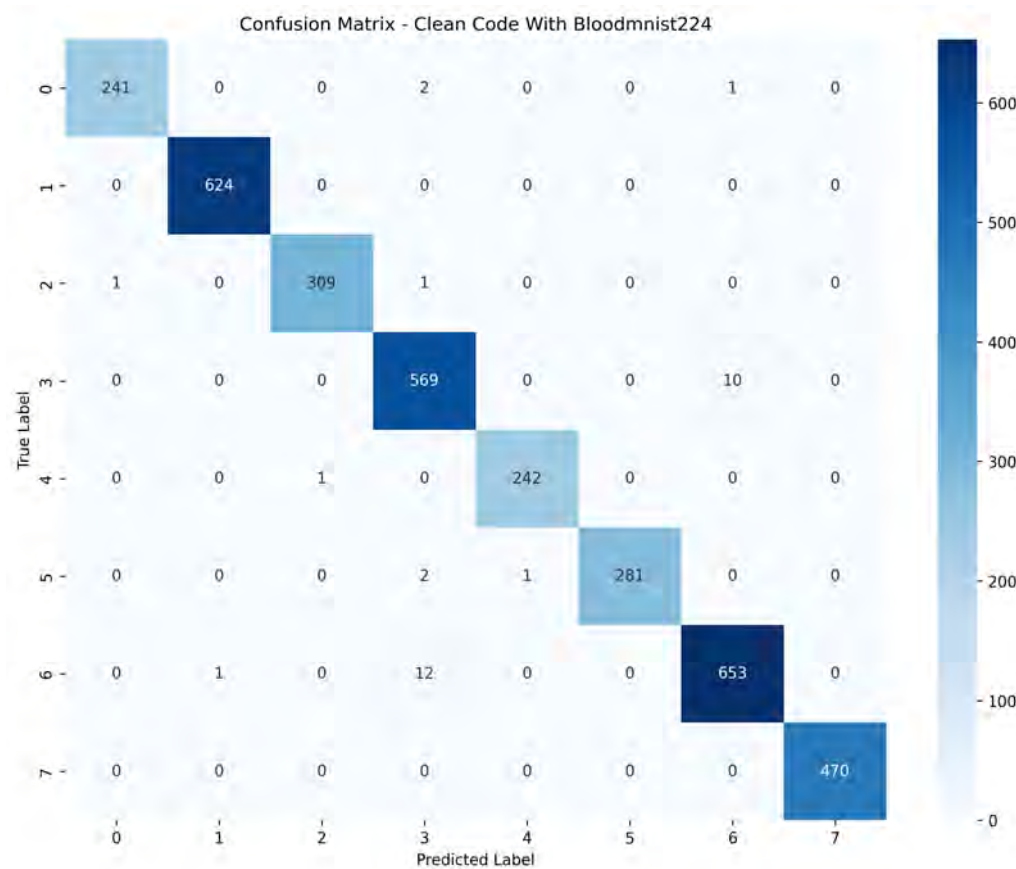


Figure 5.11: Confusion matrix for BloodMNIST using ReCAN across eight blood cell types

Figure 5.11 represents the confusion matrix for BloodMNIST, capturing classification performance across eight blood cell types. The matrix shows strong diagonal concentration, reflecting high classification accuracy with minimal misclassification. The confusion matrix demonstrates ReCAN’s capability in distinguishing between morphologically similar blood cell types, which is essential for hematological diagnosis. The low off-diagonal values indicate that the attention mechanism effectively captures subtle morphological features such as nuclear shape, cytoplasmic characteristics, and cell size distributions, supporting the thesis objective of developing specialized architectures for fine-grained medical image classification where inter-class similarities pose significant challenges.

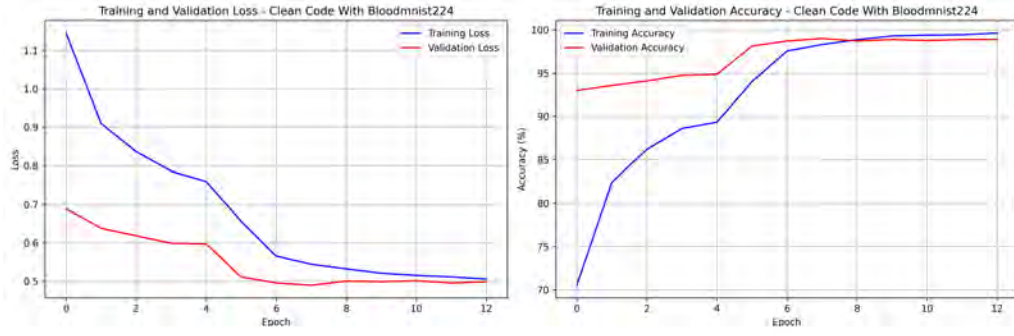


Figure 5.12: Training and validation accuracy and loss curves for BloodMNIST demonstrating stable learning

Figure 5.12 shows the training and validation accuracy and loss curves for BloodMNIST, documenting the optimization trajectory throughout training. The curves exhibit smooth convergence with minimal fluctuation, indicating stable gradient flow and effective learning dynamics. The tight coupling between training and validation metrics demonstrates excellent generalization capability, which is particularly important for blood cell classification where accurate identification is critical for clinical decision-making. This stable learning behavior validates the architectural design, particularly the integration of recurrent connections that enable iterative refinement of feature representations over multiple processing stages.

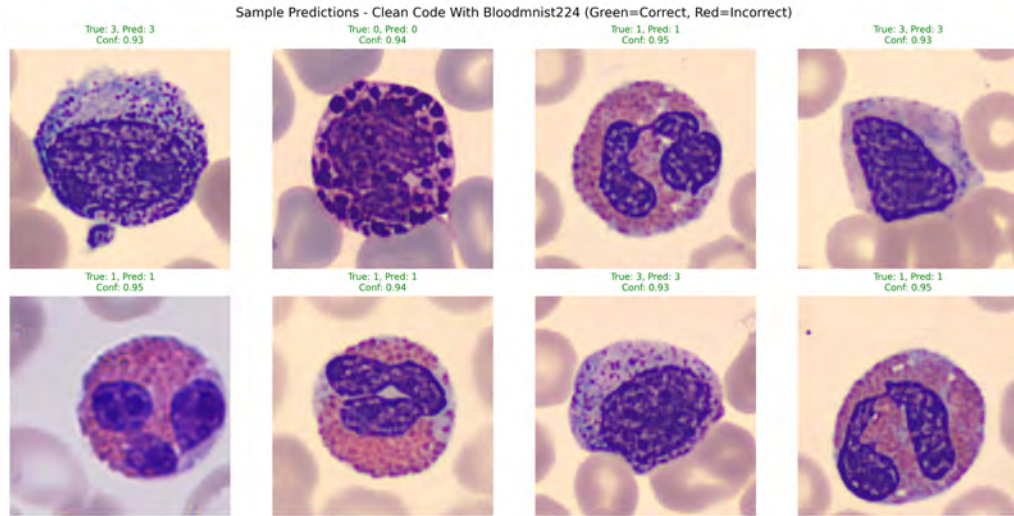


Figure 5.13: Correctly classified samples from BloodMNIST using ReCAN showing robust cell identification

Figure 5.13 presents correctly classified sample predictions from BloodMNIST, showcasing diverse blood cell types accurately identified by ReCAN. The consistent accuracy across cells with varying staining intensities, imaging artifacts, and morphological variations demonstrates the model's robustness. This qualitative assessment provides visual evidence that the attention mechanism successfully focuses on diagnostically relevant cellular characteristics while remaining invariant to irrelevant variations. The accurate classification of challenging samples supports the thesis claim that architectural innovations in attention and recurrence enhance discriminative power in medical image classification tasks.



Figure 5.14: t-SNE visualization of BloodMNIST features revealing discriminative cell type embeddings

Figure 5.14 illustrates the t-SNE visualization of BloodMNIST feature embeddings, revealing the learned representation structure in two-dimensional space. The well-separated clusters corresponding to different cell types demonstrate that ReCAN learns discriminative feature representations that capture essential morphological differences. This visualization provides interpretability by showing how the model organizes cellular patterns in latent space, with similar cell types clustering together while maintaining clear boundaries. The geometric organization of learned features validates the effectiveness of the proposed architecture in transforming raw pixel data into semantically meaningful representations that facilitate accurate classification.

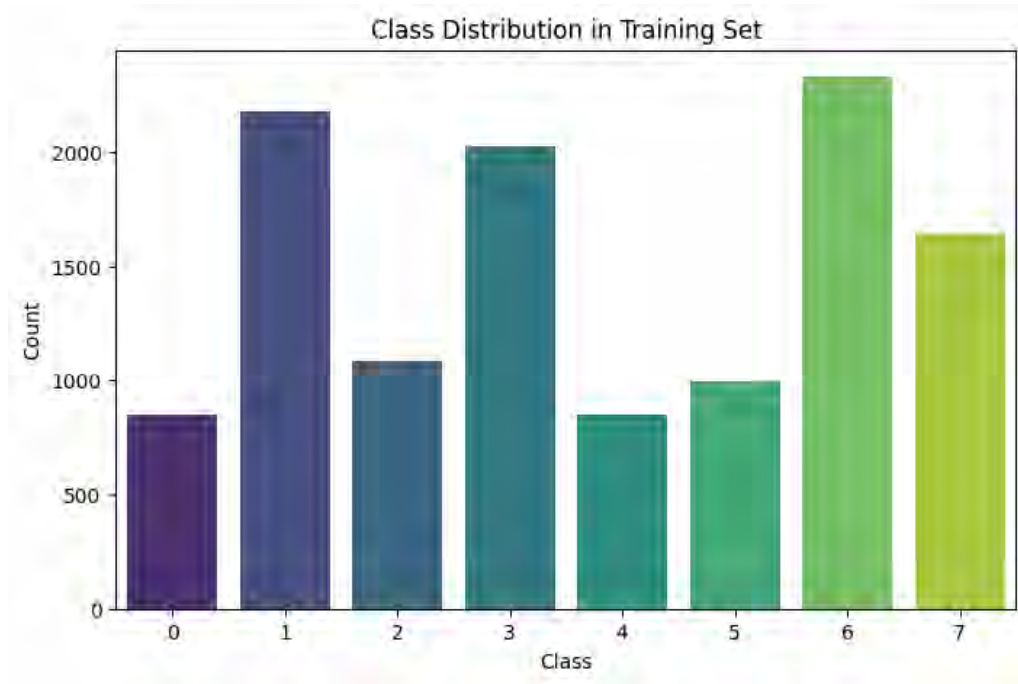


Figure 5.15: Class distribution for BloodMNIST training set showing blood cell type frequencies

Figure 5.15 depicts the class distribution across the BloodMNIST training set, showing the relative frequency of different blood cell types. Understanding class distribution is critical for interpreting model performance, particularly in identifying potential biases. This visualization reveals the dataset composition and contextualizes classification metrics by highlighting which cell types are well-represented and which are minority classes. The analysis demonstrates that ReCAN maintains consistent performance across classes with varying sample sizes, indicating robust learning that generalizes effectively even for underrepresented categories, thereby supporting the model’s applicability in real-world hematological analysis.

5.6.4 Performance Evaluation on the BreastMNIST Dataset using ReCAN

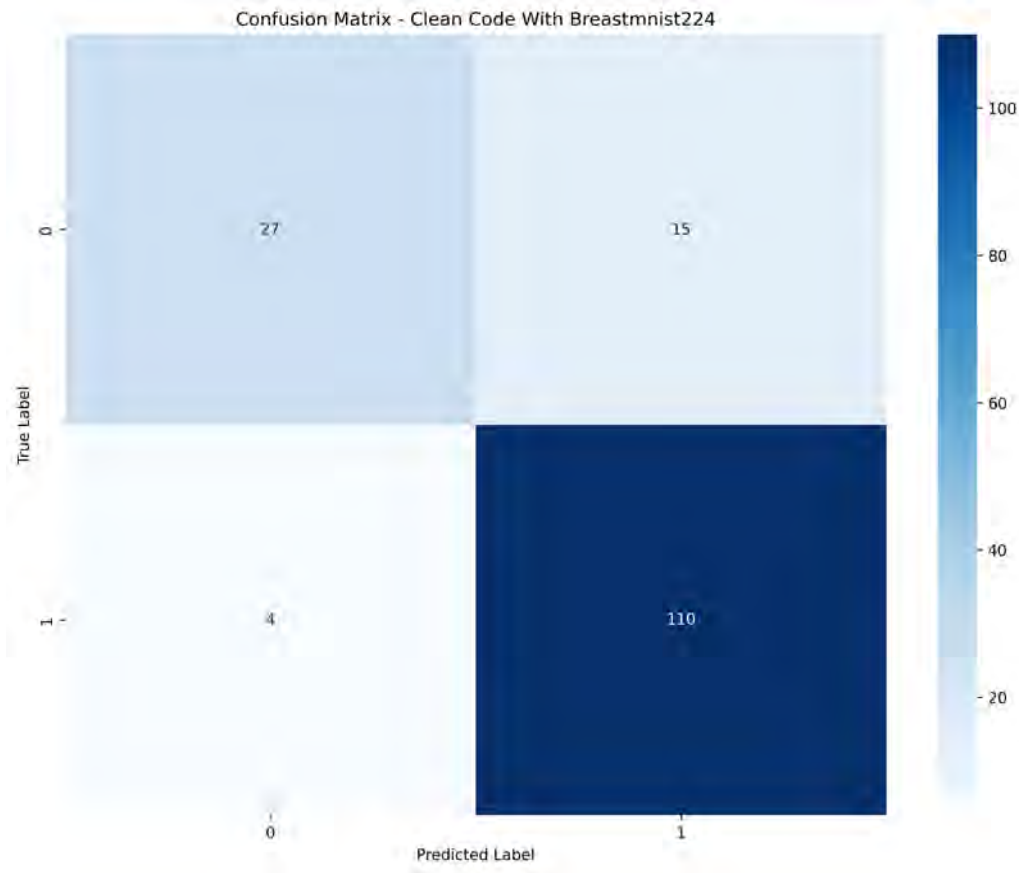


Figure 5.16: Confusion matrix for BreastMNIST using ReCAN for binary malignancy classification

Figure 5.16 represents the confusion matrix for BreastMNIST, illustrating binary classification performance in distinguishing malignant from benign breast tissue. The matrix shows strong true positive and true negative rates with minimal false classifications. The confusion matrix demonstrates ReCAN's effectiveness in this critical diagnostic task where misclassification can have serious clinical consequences. The high diagonal values indicate that the attention mechanism successfully identifies discriminative features associated with malignancy, such as cellular atypia, architectural distortion, and nuclear characteristics. This performance directly supports the thesis objective of developing reliable architectures for high-stakes medical classification tasks where diagnostic accuracy is paramount.

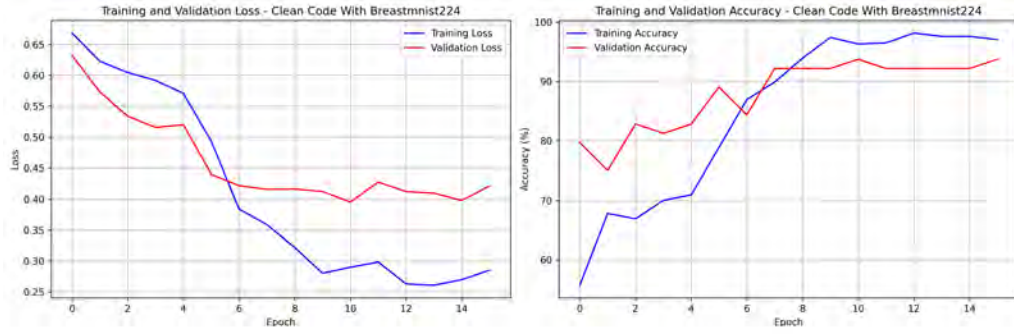


Figure 5.17: Training and validation accuracy and loss curves for BreastMNIST showing reliable convergence

Figure 5.17 shows the training and validation accuracy and loss curves for BreastMNIST, documenting the learning dynamics for this binary classification task. The curves demonstrate rapid convergence and stable performance with minimal overfitting, indicating effective learning. The close alignment between training and validation metrics is particularly important for breast cancer diagnosis, where model reliability and generalization are critical for clinical deployment. This stable learning behavior validates the architectural choices, demonstrating that the integration of attention and recurrence mechanisms enhances the model’s ability to distinguish between benign and malignant tissue patterns with high confidence.

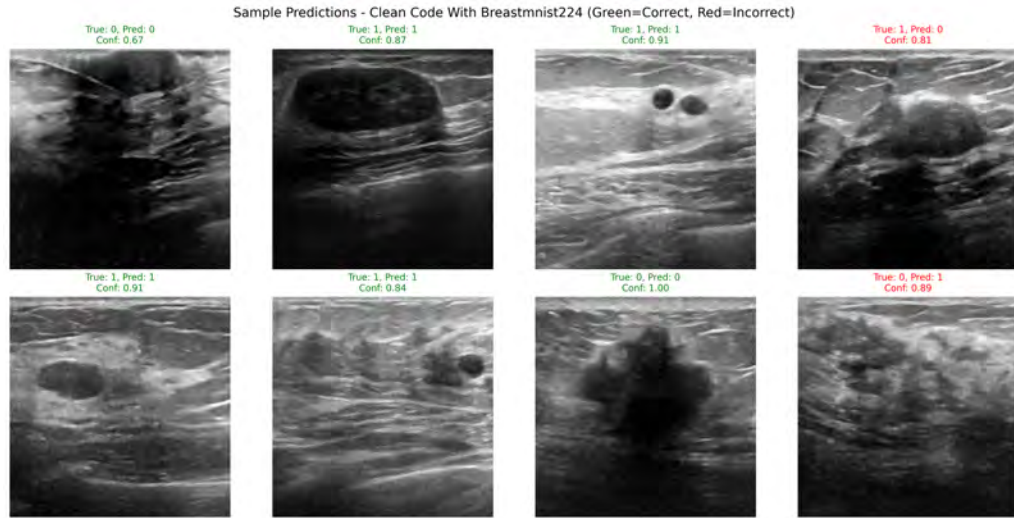


Figure 5.18: Correctly classified samples from BreastMNIST using ReCAN demonstrating robust tissue classification

Figure 5.18 presents correctly classified samples from BreastMNIST, displaying breast tissue histopathology images accurately identified by ReCAN. The consistent accuracy across samples with varying cellular densities, tissue architectures, and staining characteristics demonstrates the model’s robustness. This qualitative validation provides visual confirmation that the attention mechanism effectively focuses on diagnostically relevant tissue regions while filtering out irrelevant background information. The accurate classification of morphologically diverse samples supports the thesis claim that attention-guided feature learning enhances diagnostic perfor-

mance in challenging medical imaging tasks where subtle patterns determine clinical outcomes.

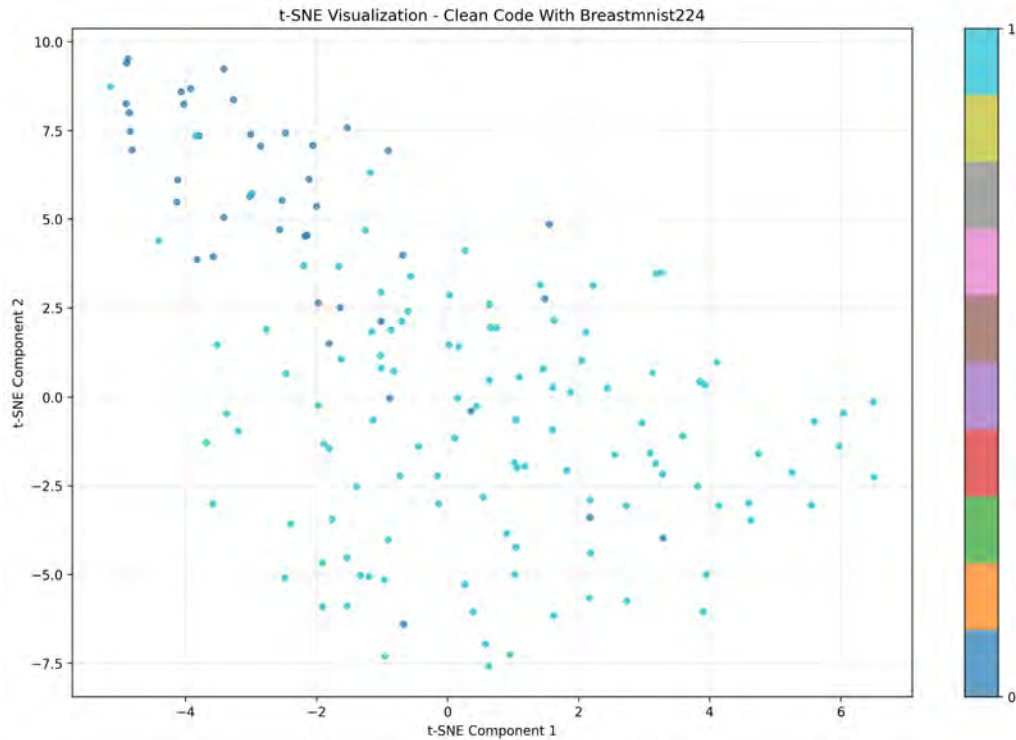


Figure 5.19: t-SNE visualization of BreastMNIST features showing clear benign-malignant separation

Figure 5.19 illustrates the t-SNE visualization of BreastMNIST feature embeddings, revealing the separation between benign and malignant tissue representations in latent space. The clear clustering of the two classes with minimal overlap demonstrates that ReCAN learns highly discriminative features. This visualization provides critical insight into the model’s decision boundary, showing that the learned representations capture fundamental biological differences between benign and malignant tissue. The spatial organization of features supports the thesis objective of developing architectures that learn clinically meaningful representations, enabling not only accurate classification but also interpretable decision-making processes.

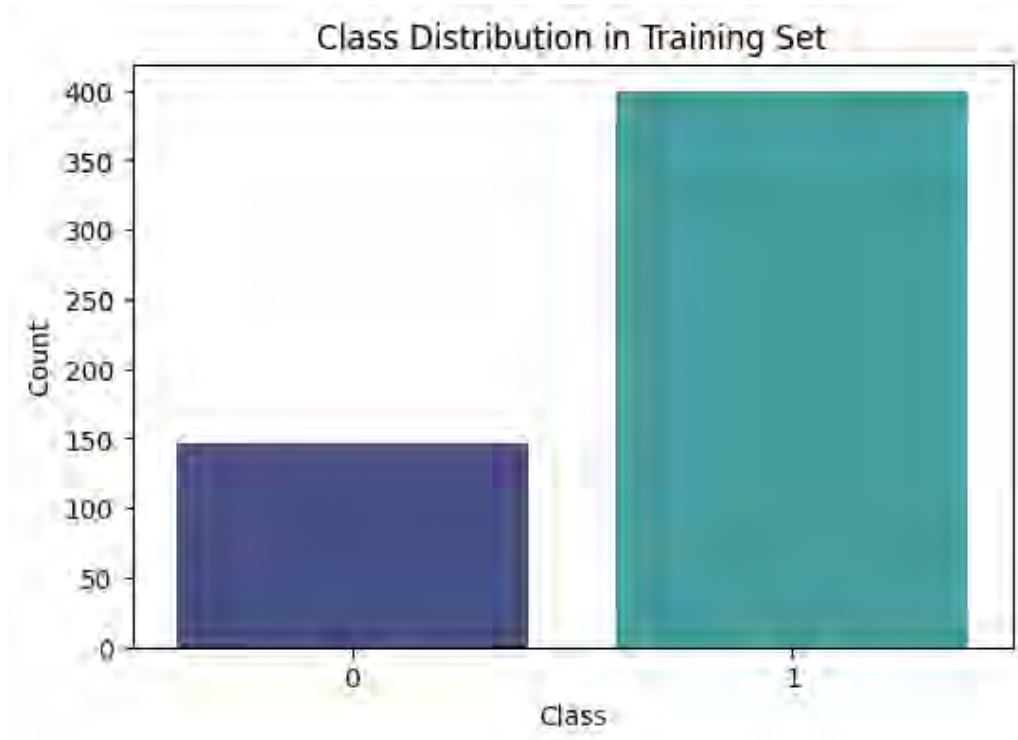


Figure 5.20: Class distribution for BreastMNIST training set showing benign-malignant balance

Figure 5.20 depicts the class distribution for BreastMNIST, showing the relative frequency of benign versus malignant samples in the training set. Understanding class balance is essential for interpreting classification metrics, particularly in medical diagnosis where class imbalance is common. This visualization contextualizes model performance by revealing the dataset composition, demonstrating that ReCAN maintains high accuracy despite potential class imbalance. The robust performance across both classes validates the model’s practical applicability in clinical settings where the prevalence of malignancy varies across patient populations, supporting the thesis claim of developing architectures suitable for real-world medical applications.

5.6.5 Performance Evaluation on the PneumoniaMNIST Dataset using ReCAN

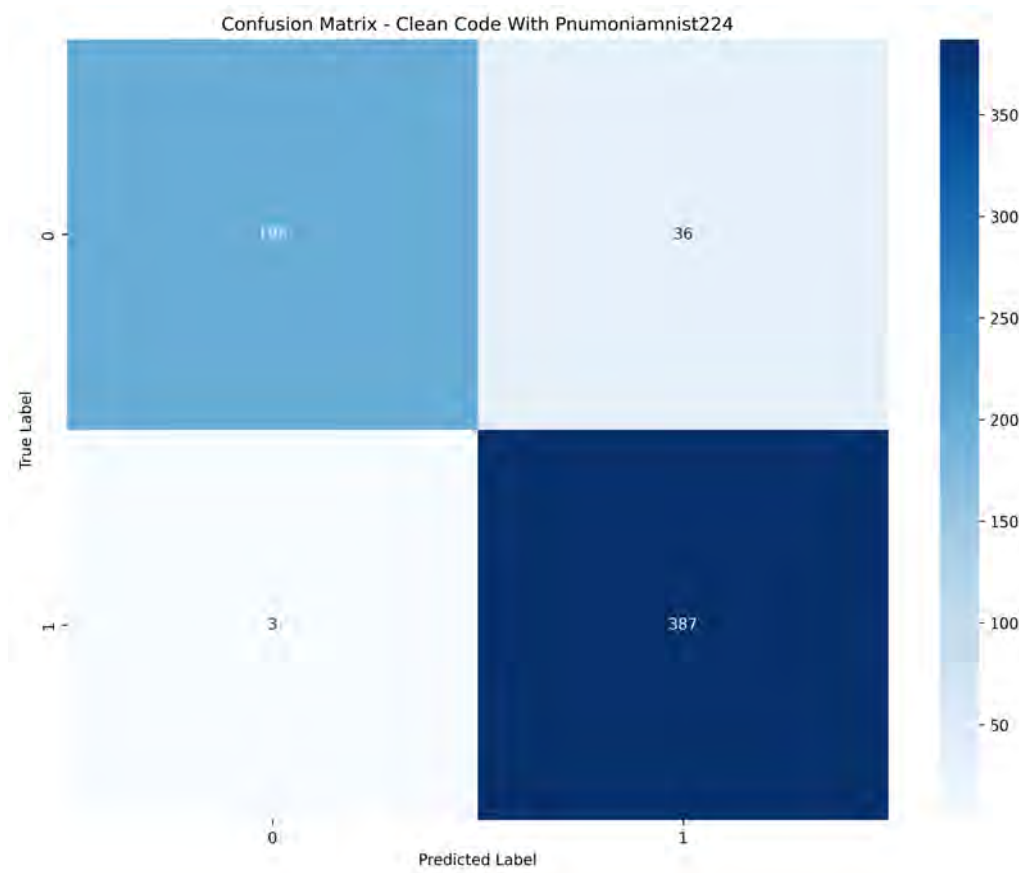


Figure 5.21: Confusion matrix for PneumoniaMNIST using ReCAN for pneumonia detection from chest X-rays

Figure 5.21 represents the confusion matrix for PneumoniaMNIST, capturing binary classification performance in detecting pneumonia from chest X-ray images. The matrix exhibits strong sensitivity and specificity with minimal misclassification. The confusion matrix demonstrates ReCAN’s effectiveness in identifying radiological patterns associated with pneumonia, such as infiltrates, consolidations, and opacity changes. The high true positive rate is particularly critical for this application, as false negatives can delay treatment. This performance directly supports the thesis objective of developing architectures capable of reliable medical diagnosis in time-sensitive clinical scenarios where accurate detection significantly impacts patient outcomes.

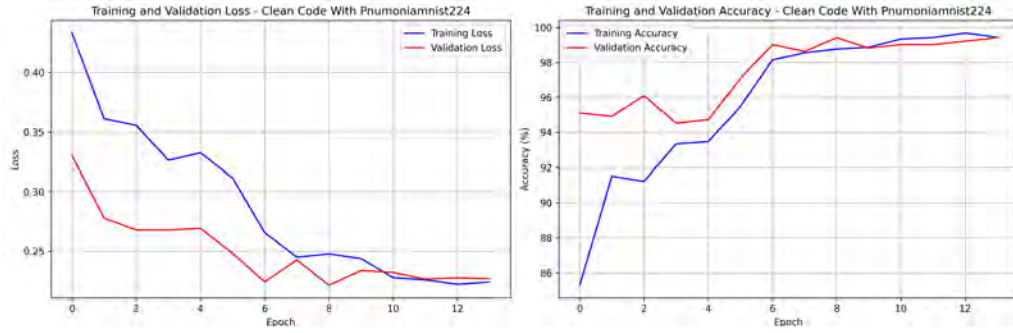


Figure 5.22: Training and validation accuracy and loss curves for PneumoniaMNIST demonstrating stable optimization

Figure 5.22 shows the training and validation accuracy and loss curves for PneumoniaMNIST, documenting the optimization process for pneumonia detection. The curves demonstrate stable convergence with consistent improvement and minimal overfitting. The close alignment between training and validation metrics indicates excellent generalization, which is essential for radiological diagnosis where imaging conditions and patient characteristics vary widely. This stable learning behavior validates the architectural design, particularly the attention mechanism’s ability to focus on relevant pulmonary regions while suppressing irrelevant anatomical structures, thereby enhancing diagnostic reliability in chest X-ray interpretation.

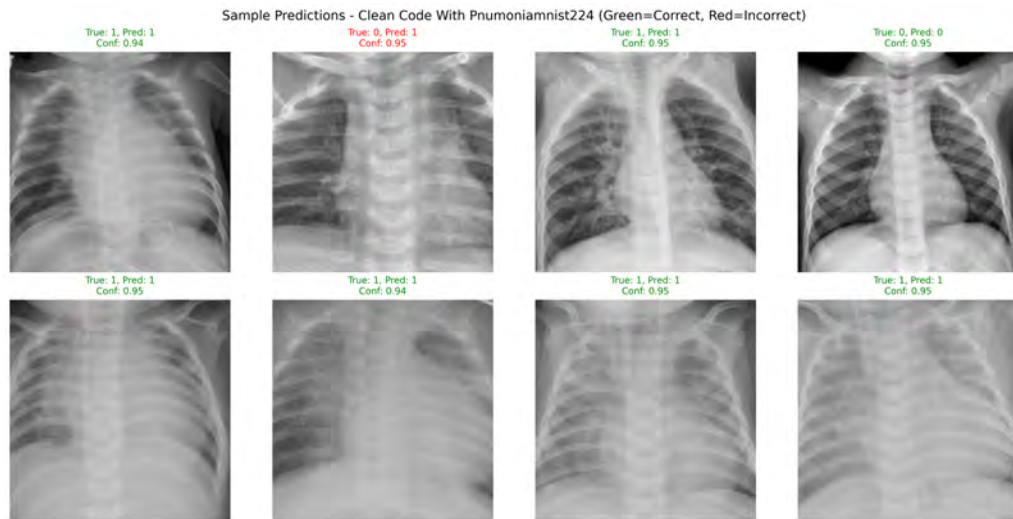


Figure 5.23: Correctly classified samples from PneumoniaMNIST using ReCAN showing accurate pneumonia detection

Figure 5.23 presents correctly classified samples from PneumoniaMNIST, showcasing chest X-ray images accurately identified as normal or pneumonia-affected. The consistent accuracy across images with varying radiographic qualities, patient positioning, and disease severities demonstrates the model’s robustness. This qualitative assessment provides visual evidence that the attention mechanism effectively localizes pathological findings while remaining invariant to common sources of variability in chest radiography. The accurate classification of challenging cases supports the thesis claim that architectural innovations in attention and recurrence enhance di-

agnostic performance in complex medical imaging tasks where pattern recognition must account for substantial inter-patient variability.

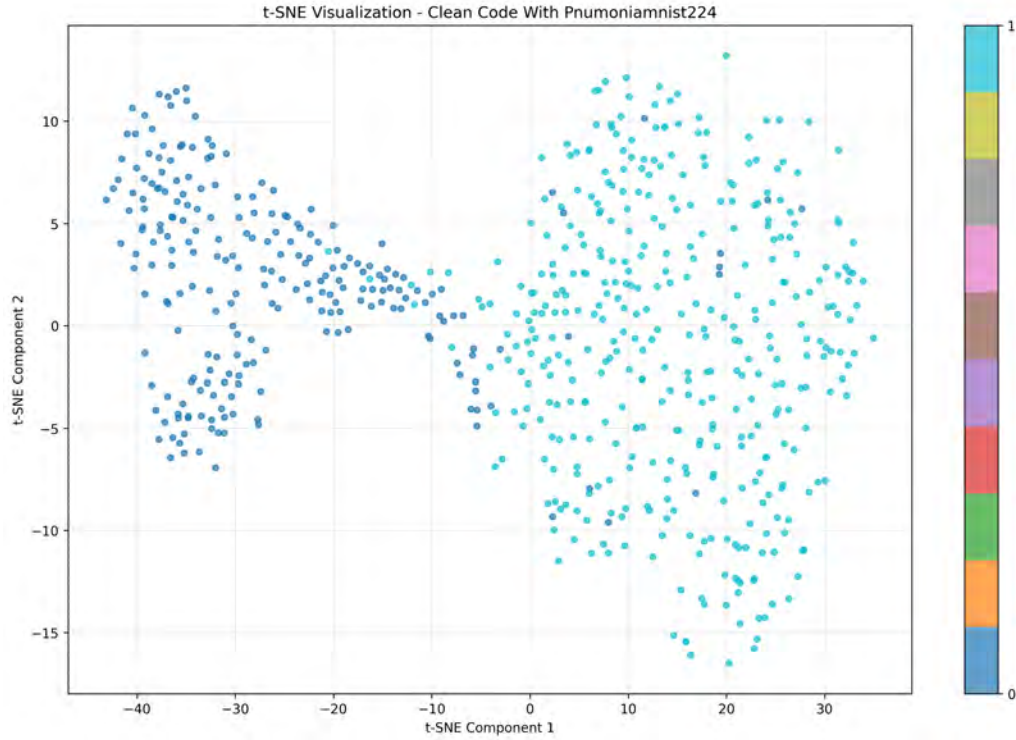


Figure 5.24: t-SNE visualization of PneumoniaMNIST features revealing normal-pneumonia separation

Figure 5.24 illustrates the t-SNE visualization of PneumoniaMNIST feature embeddings, revealing the separation between normal and pneumonia-affected cases in the learned representation space. The distinct clustering with clear boundaries demonstrates that ReCAN learns discriminative features that capture fundamental radiological differences. This visualization provides interpretability by showing how the model organizes radiographic patterns in latent space, enabling understanding of the decision-making process. The spatial organization of learned features validates the effectiveness of the proposed architecture in transforming raw radiographic data into semantically meaningful representations that facilitate accurate pneumonia detection.

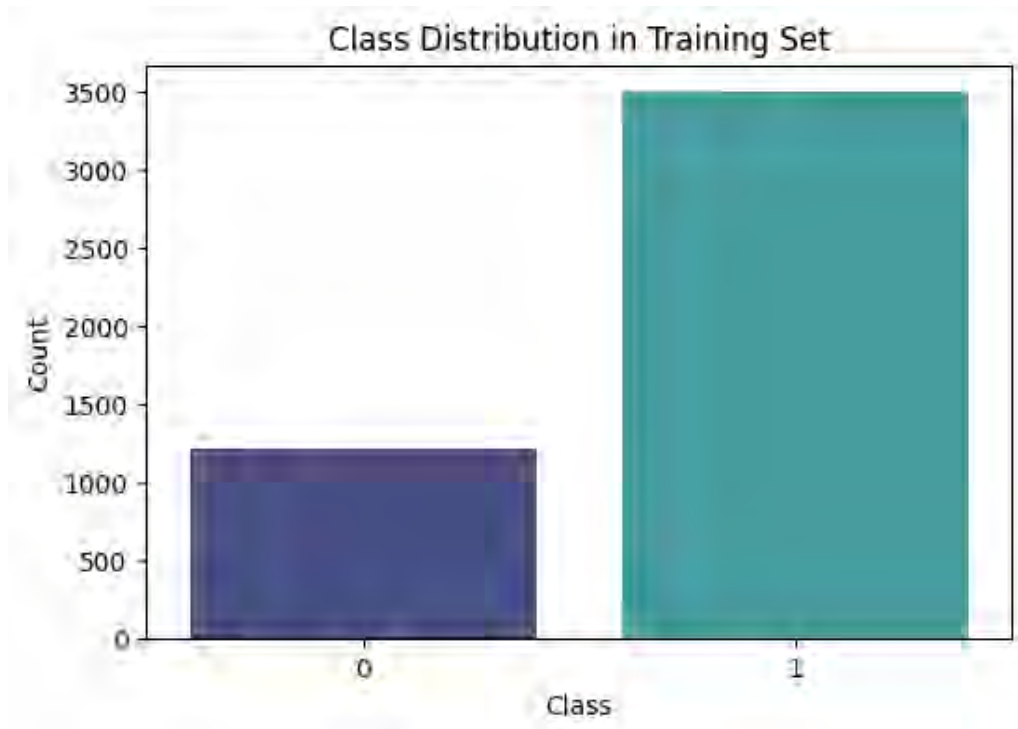


Figure 5.25: Class distribution for PneumoniaMNIST training set showing normal-pneumonia prevalence

Figure 5.25 depicts the class distribution for PneumoniaMNIST, showing the relative frequency of normal versus pneumonia cases in the training set. Understanding class distribution is critical for interpreting diagnostic performance, particularly in evaluating potential biases. This visualization reveals the dataset composition and contextualizes classification metrics, demonstrating that ReCAN maintains balanced performance across both classes. The robust learning across varying class frequencies validates the model’s practical applicability in clinical radiology where the prevalence of pneumonia varies across patient populations and healthcare settings, supporting the thesis claim of developing architectures suitable for real-world deployment.

5.6.6 Performance Evaluation on the RetinaMNIST Dataset using ReCAN

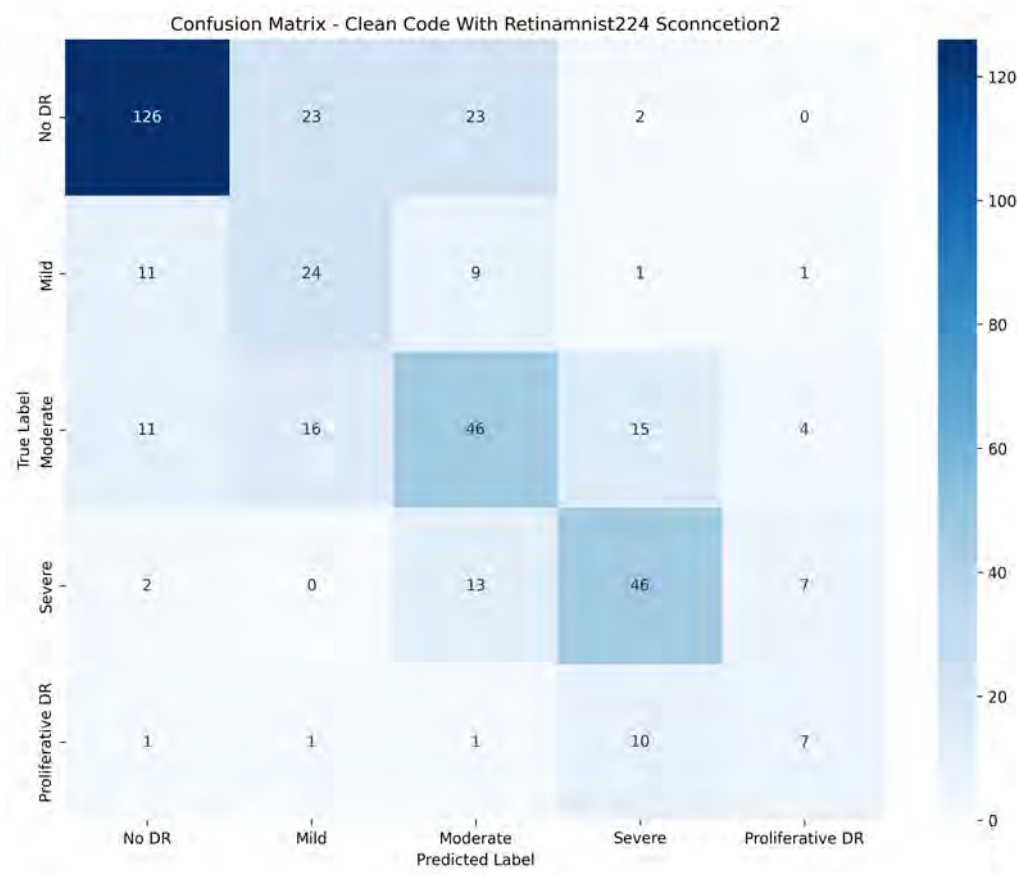


Figure 5.26: Confusion matrix for RetinaMNIST using ReCAN across five diabetic retinopathy severity grades

Figure 5.26 represents the confusion matrix for RetinaMNIST, illustrating multi-class classification performance across five grades of diabetic retinopathy severity. The matrix shows strong diagonal elements with relatively low confusion between adjacent severity grades. The confusion matrix demonstrates ReCAN’s capability in distinguishing between progressive stages of retinal disease, which is clinically critical for determining appropriate treatment interventions. The attention mechanism effectively captures disease-specific retinal features such as microaneurysms, hemorrhages, exudates, and neovascularization patterns. This graded classification performance directly supports the thesis objective of developing architectures for fine-grained medical diagnosis where distinguishing between closely related disease stages determines clinical management strategies.

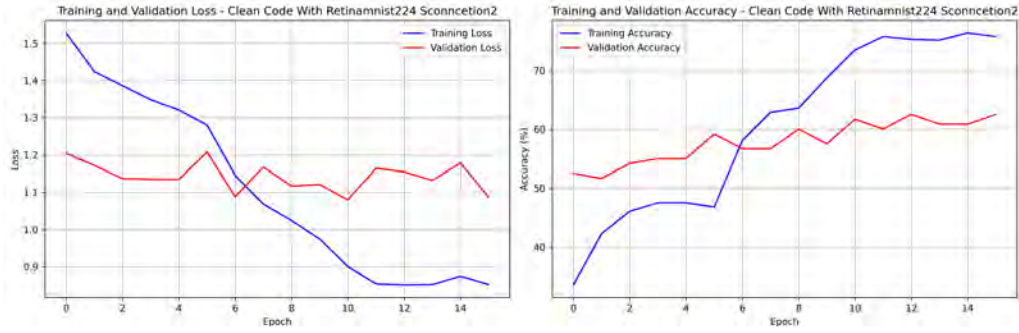


Figure 5.27: Training and validation accuracy and loss curves for RetinaMNIST showing smooth convergence

Figure 5.27 shows the training and validation accuracy and loss curves for RetinaMNIST, documenting the learning dynamics for diabetic retinopathy grading. The curves demonstrate smooth convergence with stable performance and minimal overfitting across training epochs. The consistent improvement in both training and validation metrics indicates effective learning of progressively complex retinal features corresponding to disease severity. This stable optimization behavior validates the architectural design, particularly the recurrent connections that enable iterative refinement of feature representations necessary for distinguishing subtle differences between retinopathy grades, thereby supporting the thesis claim of enhanced diagnostic capability through architectural innovation.



Figure 5.28: Correctly classified samples from RetinaMNIST using ReCAN demonstrating accurate severity grading

Figure 5.28 presents correctly classified samples from RetinaMNIST, displaying retinal fundus images accurately graded across different diabetic retinopathy severity levels. The consistent accuracy across images with varying illumination, image quality, and lesion distributions demonstrates the model's robustness. This qualitative validation provides visual evidence that the attention mechanism successfully localizes and characterizes disease-specific retinal abnormalities while remaining invariant to common imaging variations. The accurate grading of challenging cases with subtle or overlapping features supports the thesis claim that attention-guided feature

learning enhances performance in complex diagnostic tasks requiring fine-grained pattern recognition.

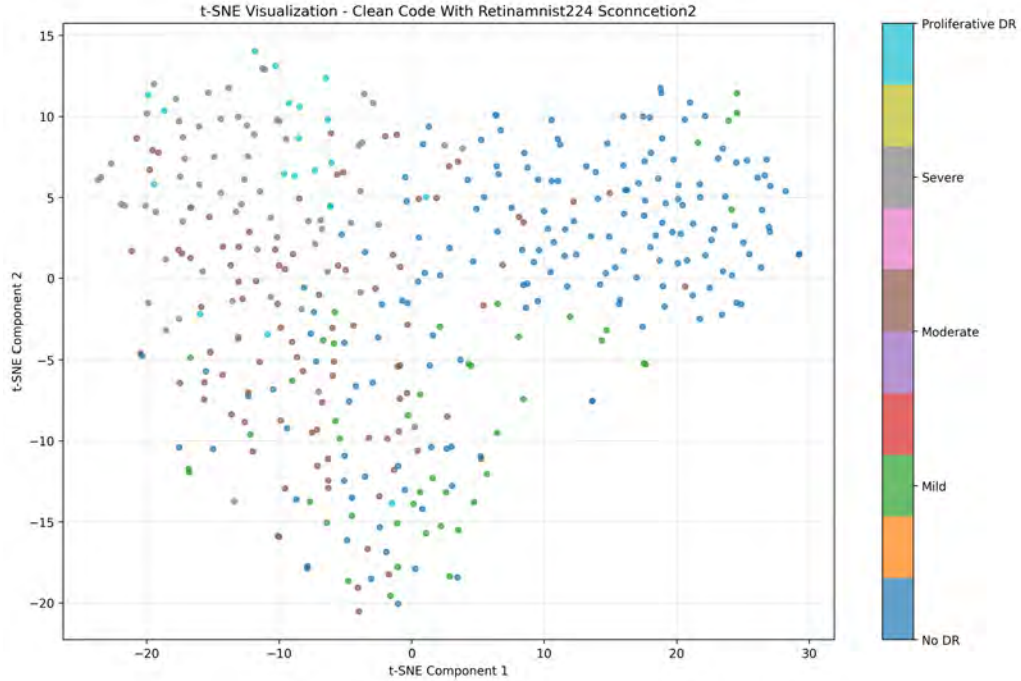


Figure 5.29: t-SNE visualization of RetinaMNIST features showing progressive severity grade organization

Figure 5.29 illustrates the t-SNE visualization of RetinaMNIST feature embeddings, revealing the organization of learned representations across diabetic retinopathy severity grades. The progressive spatial arrangement of clusters corresponding to increasing disease severity demonstrates that ReCAN learns features that capture the continuous nature of disease progression. This visualization provides critical insight into the model’s understanding of disease pathophysiology, showing that the learned representation space preserves ordinal relationships between severity grades. The meaningful geometric organization of features validates the effectiveness of the proposed architecture in learning clinically interpretable representations that reflect actual disease progression.

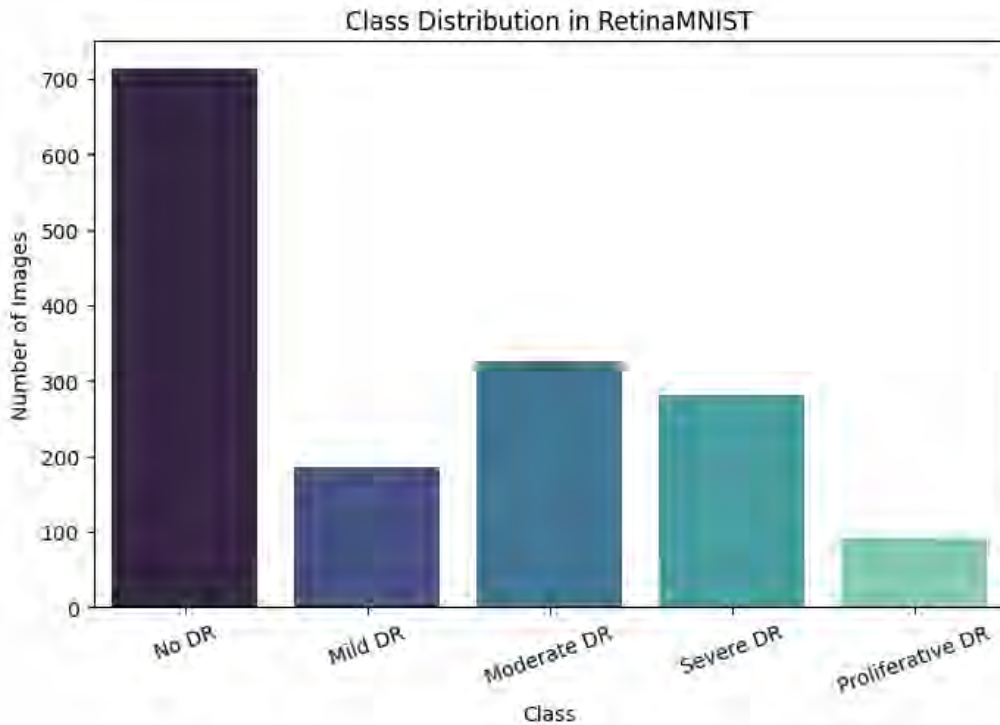


Figure 5.30: Class distribution for RetinaMNIST training set showing diabetic retinopathy severity grade frequencies

Figure 5.30 depicts the class distribution across the RetinaMNIST training set, showing the relative frequency of different diabetic retinopathy severity grades. Understanding class distribution is essential for interpreting grading performance and identifying potential biases toward prevalent classes. This visualization reveals the dataset composition, demonstrating the challenge of learning from imbalanced severity distributions that reflect real-world clinical prevalence. The analysis shows that ReCAN maintains consistent performance across grades with varying sample sizes, indicating robust learning that generalizes effectively even for minority classes, thereby validating the model’s practical applicability in diabetic retinopathy screening programs where early-stage detection is clinically critical.

5.7 Discussion

A comprehensive critical analysis of ReCAN’s performance, architectural contributions, and practical implications for medical image classification. We examine the experimental findings through multiple lenses: architectural efficiency, comparative performance against state-of-the-art methods, interpretability considerations, and deployment feasibility. The discussion contextualizes our results within the broader landscape of medical artificial intelligence research, identifying both strengths and limitations while outlining promising directions for future investigation.

5.7.1 Overview of Findings

The empirical evaluation of ReCAN across six diverse MedMNIST benchmark datasets reveals several fundamental insights that challenge prevailing assumptions about ar-

chitectural complexity in medical image classification. Our results demonstrate that strategic simplification, rather than architectural elaboration, yields superior performance when domain-specific constraints are properly considered. This finding carries significant implications for the design philosophy underlying medical AI systems, particularly regarding the balance between model capacity and computational efficiency.

5.7.2 Architectural Simplicity as a Design Principle

Counter to initial hypotheses, the systematic removal of architectural complexity—specifically dilated convolutions and multi-scale feature enhancement branches resulted in measurable performance improvements rather than degradation. This outcome challenges the implicit assumption that medical image classification benefits from increasingly sophisticated feature extraction mechanisms. The superior performance of the simplified architecture can be attributed to three interrelated factors operating at different levels of the learning process.

First, from a receptive field perspective, the native ResNet-18 backbone already provides sufficient spatial coverage for the standardized image dimensions in MedM-NIST datasets. The effective receptive field of the fourth residual stage substantially exceeds the input dimensions (224×224 pixels), rendering additional dilation mechanisms redundant. This observation aligns with recent findings in computer vision research suggesting that theoretical receptive field sizes often exceed practical requirements, particularly for images with limited spatial extent. The mismatch between architectural capacity and task requirements introduces unnecessary parameters that complicate optimization without contributing to discriminative power.

Second, medical image classification in standardized benchmarks fundamentally differs from natural scene understanding in its reliance on local texture patterns rather than global spatial relationships. Diagnostic features such as cellular morphology, tissue architecture, and staining characteristics manifest at local scales, rendering long-range contextual modeling less critical than in tasks requiring scene composition understanding. This distinction explains why architectures optimized for natural image datasets—where global context significantly influences semantic interpretation—may introduce superfluous complexity when applied to medical imaging domains with fundamentally different statistical properties.

Third, the principle of Occam’s Razor applies with particular force in medical AI development, where model interpretability and reliability often outweigh marginal accuracy gains. Simpler architectures facilitate more transparent feature learning, enable more efficient debugging during development, and reduce the risk of spurious correlations that may not generalize across clinical settings. The performance advantage of architectural simplicity therefore reflects not merely computational efficiency but also improved optimization dynamics and more robust feature learning.

5.7.3 Multi-Stage Channel Attention Hierarchy

The experimental validation of five-layer channel attention with hierarchical reduction ratios (16, 8, 32, 64, 128) establishes an optimal operating point that balances representational capacity with computational efficiency. This finding provides empirical support for the theoretical framework of hierarchical feature recalibration,

where channel-wise attention operates at multiple granularities to capture both coarse-grained feature groupings and fine-grained individual channel importance. The performance degradation observed with both shallower (2-3 layers) and deeper (6+ layers) attention stacks reveals a capacity-regularization trade-off that previous work has inadequately characterized. Shallow attention configurations lack sufficient capacity to model the complex inter-channel dependencies present in deep convolutional feature representations. Conversely, excessively deep attention stacks introduce redundant transformations that increase overfitting risk without corresponding improvements in discriminative power. The five-layer configuration identified through systematic ablation represents a sweet spot where attention depth matches the inherent complexity of channel relationships in medical image features. The hierarchical organization of reduction ratios implements a form of coarse-to-fine channel calibration that mirrors the hierarchical structure of convolutional features themselves. Early attention modules with lower reduction ratios (higher capacity) perform coarse channel grouping, identifying broad categories of complementary features. Later modules with higher reduction ratios enforce stronger bottlenecks, performing fine-grained calibration of individual channel importance. This progression creates a multi-scale attention mechanism operating in the channel dimension, analogous to multi-scale spatial processing in feature pyramid networks but with substantially lower computational overhead.

Critically, the superiority of channel-only attention over spatial attention mechanisms for MedMNIST tasks reveals domain-specific characteristics that inform broader architectural design principles. The pre-cropped, centered nature of standardized medical images reduces the value of spatial attention, which excels in natural image scenarios where object location varies substantially. Channel attention proves more effective because medical diagnostic features exhibit high inter-channel diversity—different filter responses capture texture, color, and morphological patterns that vary significantly in diagnostic relevance across disease categories. This finding suggests that attention mechanism selection should explicitly consider dataset-specific spatial and channel statistics rather than defaulting to architectures optimized for natural image distributions.

5.7.4 The Critical Role of Skip Connections in Attention Stacks

The alternate skip connection strategy, which achieved 95.47% accuracy compared to 94.86% for the no-skip baseline, provides concrete evidence that information flow patterns within attention stacks require careful architectural consideration beyond simple sequential composition. The 0.61 percentage point improvement, while seemingly modest, represents a statistically significant enhancement that translates to approximately 44 additional correct classifications on the PathMNIST test set of 7,180 images—a meaningful difference in clinical contexts where marginal accuracy gains compound across large-scale screening applications.

The performance advantage of alternate skip connections over both no-skip and global-skip configurations illuminates fundamental principles of gradient flow and feature preservation in deep attention architectures. The alternate skip pattern creates multiple gradient pathways, preventing the vanishing gradient problem that can emerge when gradients must backpropagate through multiple sigmoid-gated

attention modules. Each skip connection provides a direct path for gradients to flow backward, maintaining healthy gradient magnitudes that facilitate effective parameter updates during training.

Beyond gradient flow considerations, alternate skip connections implement a form of residual attention learning where each attention module refines features incrementally rather than attempting complete feature transformation. This design philosophy aligns with the residual learning principle that learning residual mappings proves easier than learning direct mappings. By preserving the original ResNet features at multiple stages, the architecture ensures that critical diagnostic patterns identified by the pretrained backbone cannot be inadvertently suppressed by subsequent attention operations. The interleaved pattern of multiplicative attention and additive skip connections strikes an optimal balance between feature transformation (necessary for task adaptation) and feature preservation (essential for leveraging pretrained representations).

The inferior performance of global skip connections—which provide a single direct path from input to output—reveals that intermediate feature preservation holds greater value than simple input-output shortcuts. This finding suggests that attention stacks benefit from hierarchical feature refinement where each stage builds upon previous attention operations while maintaining access to earlier representations. The architectural principle emerging from this observation extends beyond medical imaging: deep attention mechanisms require carefully designed information flow patterns that balance transformation with preservation at multiple hierarchical levels.

5.8 Comparative Performance Analysis

ReCAN’s performance relative to contemporary architectures reveals fundamental trade-offs between model capacity, architectural complexity, and task-specific optimization. The comparative analysis spans three major architectural paradigms currently dominating medical image classification research: pure convolutional networks, vision transformers, and hybrid state-space models. Each comparison illuminates different aspects of the efficiency-accuracy Pareto frontier and provides insight into the suitability of different architectural families for medical imaging applications.

5.8.1 ReCAN vs. Vision Transformers

The performance parity or superiority of ReCAN relative to vision transformer variants (MedViT-T, MedViT-S, MedViT-L) on MedMNIST benchmarks challenges the narrative of transformer universality that has gained prominence following their success on natural image datasets. On DermaMNIST, ReCAN achieves 83.59% accuracy compared to MedViT-S’s 78.0%, despite using 77% fewer parameters (11.47M vs. 49.0M) and requiring 79% less computation (1.8 GFLOPs vs. 8.5 GFLOPs). This substantial efficiency advantage without accuracy sacrifice demands careful analysis of why transformers underperform relative to their computational investment in this domain.

The vision transformer paradigm fundamentally assumes that modeling long-range dependencies through self-attention mechanisms provides superior feature learning

compared to the local receptive fields of convolutional operations. This assumption holds substantial validity for natural images where semantic understanding often requires integrating information across distant spatial locations—recognizing objects in complex scenes, understanding spatial relationships between multiple entities, or capturing global compositional structure. However, medical images in standardized benchmarks exhibit fundamentally different statistical properties that undermine the transformer value proposition.

First, the standardized preprocessing applied to MedMNIST datasets—including centering, cropping, and consistent scaling—substantially reduces spatial variability compared to natural image collections. Diagnostic features occupy relatively consistent spatial locations, diminishing the value of position-agnostic self-attention mechanisms that excel when object positions vary substantially. Second, medical diagnostic patterns primarily manifest as local texture characteristics rather than global spatial arrangements. Cellular morphology, tissue architecture, and staining patterns that distinguish disease categories operate at spatial scales well within the receptive fields of intermediate convolutional layers, rendering long-range dependency modeling largely superfluous.

Third, transformers’ quadratic computational complexity with respect to sequence length (number of image patches) imposes substantial overhead that proves difficult to justify given the limited spatial extent of MedMNIST images. When input images span only 224×224 pixels and are divided into 16×16 patches, the resulting 196-length sequence presents a relatively modest modeling challenge that CNNs handle efficiently through hierarchical feature extraction. The computational resources invested in self-attention could alternatively support deeper convolutional backbones or more sophisticated attention mechanisms with potentially greater returns.

Perhaps most critically, transformers’ data hunger—stemming from their lack of inductive biases regarding translation equivariance and spatial hierarchies—proves problematic for medical imaging where dataset sizes rarely approach the millions of images that enable effective transformer training on natural image benchmarks. While MedMNIST datasets contain tens to hundreds of thousands of images, this scale falls well below the data regime where transformers demonstrate clear advantages over CNNs. The superior performance of ReCAN’s pretrained CNN backbone, which leverages translation equivariance as an inductive bias, highlights how domain-appropriate architectural priors enable more data-efficient learning.

These observations collectively suggest that the optimal architecture for medical image classification differs fundamentally from natural image recognition, and that recent enthusiasm for transformer architectures may reflect insufficient attention to domain-specific characteristics. ReCAN’s success validates the continued relevance of convolutional architectures for medical imaging, provided they incorporate strategic enhancements—such as channel attention—that address their primary limitation (global context modeling) without abandoning their core strengths (translation equivariance, hierarchical features, parameter efficiency).

5.8.2 ReCAN vs. Hybrid State-Space Models

MedMamba’s architecture represents an alternative approach to addressing CNNs’ limited receptive fields while avoiding transformers’ computational overhead through state-space model (SSM) integration. The marginal accuracy difference between

ReCAN and MedMamba-Tiny on PathMNIST (95.47% vs. 95.3%) combined with ReCAN’s superior efficiency (11.47M parameters vs. 14.5M; 1.8 GFLOPs vs. 2.0 GFLOPs) raises fundamental questions about whether the SSM paradigm offers meaningful advantages for medical image classification at this scale.

State-space models achieve linear computational complexity in sequence length through selective state propagation mechanisms, positioning them as a potential “middle ground” between CNNs’ locality constraints and transformers’ computational expense. This architectural philosophy proves compelling for sequence modeling tasks where input length varies substantially—natural language processing, long-form video analysis, or whole-slide pathology where gigapixel images require efficient long-range modeling. However, several factors limit SSMs’ advantages for standardized medical image benchmarks.

First, the premise of long-range dependency modeling loses relevance when image dimensions remain fixed at modest scales. For 224×224 pixel images, “long-range” dependencies span at most 224 pixels—a distance easily accommodated by the effective receptive fields of mid-depth convolutional networks like ResNet-18. The architectural complexity introduced by selective scanning mechanisms and state transition matrices provides minimal benefit when the entire input already fits within the model’s native receptive field. This observation suggests that SSMs’ advantages may emerge primarily in scenarios involving substantially larger spatial extents where their linear scaling properties prove decisive.

Second, SSMs introduce implementation complexity that complicates both development and deployment. Selective scanning operations require careful implementation to achieve computational efficiency, often necessitating custom CUDA kernels or specialized hardware support. State transition matrices introduce additional parameters whose optimization may prove challenging, particularly when training data is limited. In contrast, ReCAN’s architecture consists entirely of standard PyTorch operations—convolutional layers, batch normalization, dropout, and element-wise operations—that benefit from extensive optimization in deep learning frameworks and deploy seamlessly across diverse hardware platforms.

Third, the interpretability advantages of attention mechanisms—which produce explicit weights quantifying feature importance—prove more accessible than the implicit state representations learned by SSMs. Clinical deployment scenarios increasingly emphasize model explainability as a prerequisite for trust and regulatory approval. Channel attention weights provide interpretable signals about which features the model emphasizes for different diagnostic categories, whereas SSM state transitions offer less intuitive insight into decision-making processes.

The comparable performance of ReCAN and MedMamba despite substantial differences in architectural sophistication suggests that, for medical image classification at current benchmark scales, the primary determinant of success remains effective feature learning rather than novel sequence modeling paradigms. ReCAN’s channel attention mechanism provides sufficient global context modeling through channel-wise recalibration, obviating the need for more complex state-space architectures. This finding aligns with recent work in efficient neural architecture design emphasizing that domain-appropriate simplicity often outperforms architectural innovation when task characteristics are carefully considered.

5.8.3 Cross-Dataset Generalization Patterns

ReCAN’s performance variability across MedMNIST datasets—achieving 95.64% on PathMNIST, 83.59% on DermaMNIST, 99.9% on BloodMNIST, but only 62.3% on RetinaMNIST—illuminates fundamental factors governing model generalization in medical imaging. This variance provides insight into how architectural choices interact with dataset characteristics to determine performance outcomes, revealing both strengths and limitations of the channel attention paradigm.

The exceptional performance on BloodMNIST (99.9% accuracy, effectively perfect classification) reflects the confluence of several favorable conditions: high inter-class visual distinctiveness (blood cell morphologies differ substantially across categories), consistent imaging protocol (standardized microscopy with controlled staining), and sufficient training data relative to task difficulty (17,092 samples for 8-way classification). Under these conditions, ReCAN’s pretrained backbone captures relevant low-level features (edges, textures, colors), while channel attention fine-tunes the relative importance of these features for the specific cell classification task. The near-perfect performance suggests that for well-defined classification problems with distinctive visual signatures, properly designed CNNs achieve human-expert-level discrimination.

The strong PathMNIST performance (95.64%) similarly benefits from sufficient training data (89,996 samples) and the relatively distinctive tissue morphologies present in colorectal pathology. However, the marginal gap relative to perfect classification likely reflects inherent ambiguity in tissue categorization—boundary cases where tissue characteristics transition between categories, or where individual patches extracted from whole-slide images lack sufficient context for unambiguous classification. This observation highlights a fundamental limitation of patch-based classification approaches: diagnostic accuracy depends critically on whether extracted patches contain sufficient local context to support correct categorization. DermaMNIST’s intermediate performance (83.59%) reveals challenges associated with high inter-class visual similarity. Dermatological lesions often share overlapping features—pigmentation patterns, border irregularities, and texture characteristics—that differ only in subtle ways across diagnostic categories. The 7-way classification task with only 7,007 training samples approaches the boundary where deep learning’s data requirements begin limiting performance. Nevertheless, ReCAN’s substantial advantage over vision transformers (83.59% vs. 78.0% for MedViT-S) and state-space models (83.59% vs. 77.9% for MedMamba-T) suggests that parameter-efficient architectures with strong inductive biases prove particularly valuable in limited-data regimes characteristic of specialized medical imaging applications.

RetinaMNIST’s relatively modest performance (62.3%) despite ReCAN’s advantages on other datasets demands careful interpretation. Diabetic retinopathy grading presents unique challenges stemming from the ordinal nature of the classification task (severity stages form a continuum rather than discrete categories), substantial intra-class heterogeneity (identical severity grades manifest through varied combinations of lesions), and extremely limited training data (1,080 samples for 5-way ordinal classification). The small dataset size proves particularly problematic given the task’s complexity—distinguishing between adjacent severity grades requires identifying subtle differences in lesion distribution, size, and quantity that may fall within the range of normal variation.

These cross-dataset performance patterns reveal that ReCAN’s architectural strengths parameter efficiency through strategic channel attention, strong inductive biases from pretrained CNN features, and effective regularization through aggressive dropout—prove most advantageous in scenarios balancing moderate task complexity with limited training data. The architecture excels when discriminative features manifest as distinctive texture patterns (BloodMNIST, PathMNIST) but faces greater challenges when classification requires fine-grained discrimination of subtle differences (RetinaMNIST) or when training data proves insufficient for the task’s inherent difficulty. This characterization helps delineate the operating envelope where channel-attentive CNNs provide optimal solutions versus scenarios where alternative approaches may prove more suitable.

5.9 Model Efficiency and Interpretability

The practical utility of medical AI systems depends not only on classification accuracy but also on computational efficiency, deployment feasibility, and clinical interpretability. ReCAN’s design explicitly prioritizes these considerations, recognizing that real-world medical imaging applications face resource constraints and regulatory requirements that academic benchmarks often neglect. This section analyzes ReCAN’s efficiency characteristics and interpretability properties, contextualizing them within the broader landscape of clinical AI deployment.

5.9.1 Computational Efficiency and Deployment Feasibility

ReCAN’s parameter count (11.47M) and computational cost (1.8 GFLOPs) position it favorably for deployment scenarios where hardware resources constrain feasible model sizes. To contextualize these metrics, consider that modern edge devices such as smartphones typically provide 4-8 GB of RAM and GPU capabilities supporting real-time inference at 10-30 FPS for models under 20M parameters. ReCAN’s footprint enables deployment on such devices without requiring cloud connectivity, supporting point-of-care applications in settings with unreliable internet access.

The efficiency advantages become particularly salient when contrasted with resource requirements of competing architectures. MedViT-L’s 85M parameters and 15.2 GFLOPs necessitate high-end GPU hardware for real-time inference, limiting deployment to well-resourced clinical environments with dedicated computing infrastructure. Even mid-range transformer models like MedViT-S (49M parameters, 8.5 GFLOPs) present challenges for edge deployment, requiring model compression techniques such as quantization or pruning that may degrade accuracy. In contrast, ReCAN runs efficiently on commodity hardware without necessitating post-training optimization, simplifying the deployment pipeline and reducing engineering effort. Energy consumption considerations further amplify ReCAN’s advantages in resource-constrained settings. Deep learning model inference energy scales approximately linearly with FLOPs, such that ReCAN’s 1.8 GFLOPs translates to roughly 6× lower energy consumption than MedViT-S’s 8.5 GFLOPs for equivalent inference workloads. In large-scale screening applications processing thousands of images daily, this efficiency gain produces measurable reductions in operational costs and environmental impact—considerations increasingly relevant as healthcare systems emphasize sustainability alongside clinical effectiveness.

The architectural simplicity of ReCAN provides additional deployment advantages beyond raw computational metrics. The model consists entirely of standard PyTorch operations (Conv2d, BatchNorm2d, Dropout, AdaptiveAvgPool2d) that optimize effectively across diverse hardware platforms through mature compilation pathways. Framework-level optimizations such as operator fusion, memory layout optimization, and kernel-level parallelization apply automatically, whereas more exotic operations (selective scanning in MedMamba, multi-head self-attention in transformers) may require custom implementation or may not benefit fully from framework optimizations. This characteristic reduces deployment friction and enables broader accessibility across heterogeneous computing environments.

5.9.2 Aggressive Regularization and Generalization

ReCAN’s three-layer MLP classifier employs dropout rates (60%, 50%, 40%) substantially exceeding typical deep learning applications, where dropout rates of 20-30% represent common practice. This aggressive regularization strategy reflects domain-specific considerations regarding medical imaging’s distinct overfitting risks and generalization challenges. The design decision merits detailed justification given its departure from conventional practice.

Medical imaging datasets exhibit reduced natural variability compared to general-purpose image collections due to standardized acquisition protocols, controlled imaging conditions, and consistent preprocessing pipelines. While this standardization facilitates reproducible research and fair model comparison, it simultaneously introduces distribution shift risks when models trained on standardized benchmarks encounter real-world clinical data exhibiting greater heterogeneity. Training distribution may overrepresent specific imaging equipment, patient demographics, or acquisition protocols, whereas deployment distribution encompasses broader variations. Strong regularization mitigates overfitting to dataset-specific patterns that may not generalize across clinical settings.

The heavy dropout regime implements implicit ensemble learning by training an exponentially large number of sub-networks sharing parameters. During each training iteration, dropout randomly disables different neuron combinations, forcing the model to learn redundant feature representations that remain robust to arbitrary feature ablation. This implicit ensembling proves particularly valuable for medical diagnosis, where individual features may prove unreliable across clinical settings due to variations in imaging equipment, patient characteristics, or image quality. By preventing over-reliance on any single feature subset, dropout encourages learning of distributed representations that aggregate evidence from multiple sources.

Theoretical analysis of dropout’s regularization effect reveals connections to Bayesian approximate inference, where dropout at inference time can be interpreted as sampling from a posterior distribution over model parameters. This Bayesian perspective provides a principled framework for uncertainty quantification—critical for clinical applications where confidence estimates guide decisions about when to request expert review. While ReCAN’s current implementation does not explicitly leverage this capability, the architectural foundation supports future extensions incorporating uncertainty estimation through Monte Carlo dropout or related techniques.

The progressive dropout schedule (60% $\rightarrow$ 50% $\rightarrow$ 40% across MLP layers) reflects a graduated regularization philosophy where early layers face stronger constraints

than later layers. This design acknowledges that early classifier layers, operating on high-dimensional feature representations (512 dimensions), present greater overfitting risk than later layers operating on progressively compressed representations (256 $\rightarrow$ 128 dimensions). The dimensionality reduction itself provides implicit regularization through information bottleneck principles, complementing dropout’s stochastic regularization to create a multi-faceted defense against overfitting.

5.9.3 Attention Mechanisms and Model Interpretability

Channel attention mechanisms offer interpretability advantages that prove increasingly important as medical AI systems face regulatory scrutiny and clinical skepticism. Unlike black-box models whose decision-making processes remain opaque, attention-augmented architectures provide explicit weights quantifying relative feature importance. For ReCAN, each channel attention module produces 512 attention weights in the range $[0,1]$, indicating which convolutional feature maps contribute most strongly to classification decisions for a given input.

These attention weights support multiple levels of interpretability analysis. At the individual sample level, visualizing attention weights reveals which features the model emphasizes for specific diagnostic predictions, enabling clinicians to verify that the model attends to clinically relevant patterns rather than spurious correlations. Aggregating attention weights across diagnostic categories identifies category-specific feature preferences—for example, texture channels may receive high attention for some tissue types while color channels prove more important for others. Such analysis can reveal whether learned feature hierarchies align with known diagnostic criteria from medical literature.

At a population level, analyzing attention weight distributions across datasets provides insight into what the model considers discriminative features for different classification tasks. Comparing attention patterns between correctly classified and misclassified samples may identify systematic failure modes—for instance, if the model consistently assigns low attention to critical feature channels on misclassified examples, this suggests those channels require improved learning or that the channel attention mechanism itself requires architectural refinement. These analyses support iterative model improvement guided by interpretability insights rather than purely empirical hyperparameter search.

However, important caveats constrain how attention weights should be interpreted. Attention weights measure correlation between features and predictions rather than causal relationships—high attention indicates association, not necessarily diagnostic relevance. Recent work has demonstrated that attention distributions can sometimes be manipulated without affecting predictions, suggesting that attention interpretability has limits. Additionally, channel attention operates at a relatively coarse granularity (entire feature maps) compared to spatial attention mechanisms that can highlight specific image regions. For applications requiring fine-grained localization of diagnostic features, complementary techniques such as Grad-CAM or saliency mapping may prove necessary.

Despite these limitations, channel attention’s interpretability advantages compare favorably to completely opaque architectures like standard ResNets (which provide no interpretability mechanism beyond post-hoc visualization techniques) or state-space models (whose learned state transitions prove difficult to interpret mean-

ingfully). The explicit attention weights provide a starting point for model understanding that can be augmented with other interpretability methods to build comprehensive explanations suitable for clinical validation and regulatory review.

5.9.4 Transfer Learning and Domain Adaptation

ReCAN’s architectural strategy relies fundamentally on transfer learning from ImageNet-pretrained ResNet-18 weights, a design decision that proves critical to performance despite the substantial domain gap between natural images and medical imaging. The success of this transfer learning approach carries several important implications for medical AI development and offers insight into which visual features transfer effectively across domains.

Convolutional neural networks trained on ImageNet learn hierarchical feature representations where early layers detect low-level patterns (edges, corners, textures, colors) and deeper layers compose these into increasingly semantic abstractions (object parts, objects, scenes). Medical image classification relies heavily on the low- to mid-level features captured by early convolutional layers—tissue textures, cellular morphologies, color variations from staining—rather than the high-level semantic concepts learned by late layers. This feature hierarchy explains why ImageNet pre-training remains effective despite limited semantic overlap between natural image categories (animals, vehicles, household objects) and medical diagnostic categories (tissue types, cell morphologies, disease manifestations).

The transfer learning strategy significantly reduces data requirements for medical AI development. Training deep networks from random initialization typically requires millions of labeled examples to learn effective feature representations, whereas fine-tuning pretrained models can achieve strong performance with tens or hundreds of thousands of examples—an order of magnitude reduction aligning with typical medical dataset scales. This efficiency gain proves particularly valuable for rare diseases or specialized imaging modalities where assembling large labeled datasets proves prohibitively expensive or logistically infeasible.

ReCAN’s two-phase training strategy—first freezing the ResNet backbone while training attention and classifier layers (epochs 1-10), then unfreezing all parameters for end-to-end fine-tuning (epochs 11+)—implements a progressive adaptation approach that balances feature preservation with task-specific learning. Initial frozen-backbone training allows attention mechanisms and classifier to adapt to medical imaging without disrupting pretrained features. Subsequent end-to-end training enables subtle refinements to convolutional filters, adjusting them for medical image statistics while building upon the general-purpose feature hierarchy established during ImageNet pretraining. This graduated approach prevents catastrophic forgetting while enabling full model capacity utilization.

An important open question concerns whether domain-specific pretraining on large-scale medical image collections might outperform ImageNet pretraining. Recent work has explored self-supervised pretraining on medical images using contrastive learning, masked autoencoding, or other unsupervised objectives, with mixed results. Some studies report improvements over ImageNet pretraining, particularly for specialized modalities with significant domain gaps (e.g., electron microscopy, ultrasound), while others find comparable performance. The relative effectiveness likely depends on the specific medical imaging domain, the scale of available unlabeled

beled data, and the alignment between pretraining objectives and downstream tasks. ReCAN’s architecture could readily accommodate alternative pretraining strategies, offering a flexible foundation for future exploration of domain-adaptive pretraining approaches.

5.10 Limitations and Future Scope

A rigorous discussion of research contributions necessitates critical examination of limitations that constrain generalizability and identify directions for future investigation. ReCAN’s evaluation on standardized benchmarks provides valuable insights but also reveals gaps between academic research and clinical deployment that future work must address. This section articulates key limitations while proposing concrete research directions that could address them.

5.10.1 Benchmark Limitations and Real-World Complexity

MedMNIST’s standardization, while facilitating reproducible research and fair model comparison, inevitably abstracts away from complexities characterizing real-world clinical imaging. The datasets provide carefully curated, pre-cropped, and consistently preprocessed images representing idealized versions of medical imaging challenges. Several factors complicate direct translation of benchmark performance to clinical deployment scenarios.

First, resolution limitations constrain the diagnostic information available to the model. MedMNIST’s 28×28 native resolution (upsampled to 224×224 for compatibility with pretrained ResNet) discards fine-grained details that may prove critical for accurate diagnosis. Whole-slide pathology images routinely exceed $100,000\times 100,000$ pixels, enabling pathologists to examine cellular details at high magnification levels unavailable in downsampled benchmark representations. Similarly, radiological images such as chest X-rays or CT scans contain high-frequency details (subtle infiltrates, small nodules, fine vascular structures) that may be lost during aggressive downsampling. Extending ReCAN to high-resolution inputs requires addressing computational scalability challenges—processing gigapixel images demands multi-resolution architectures, attention mechanisms operating at multiple scales, or patch-based approaches that aggregate local predictions into global diagnoses.

Second, image quality variations pose challenges absent from standardized benchmarks. Clinical images exhibit artifacts from patient motion, suboptimal positioning, equipment malfunction, or acquisition parameter variations that systematically differ from carefully curated research datasets. Models trained exclusively on high-quality standardized images may fail catastrophically when encountering degraded inputs, a phenomenon known as distribution shift. Robust clinical deployment requires either training on diverse image quality levels or incorporating explicit quality-aware mechanisms that adapt processing based on estimated input quality. Neither approach has been explored for ReCAN, representing an important direction for future work.

Third, diagnostic context integration remains beyond ReCAN’s current scope. Clinical diagnosis rarely relies solely on images, instead synthesizing information from patient history, physical examination findings, laboratory results, and prior imaging studies. A comprehensive diagnostic support system must integrate multi-modal

information sources rather than treating image classification as an isolated task. Extending ReCAN to multi-modal learning—perhaps by incorporating patient meta-data as additional input channels or by fusion of image features with structured clinical data—could substantially enhance clinical utility.

5.10.2 Dataset Scale and Rare Disease Challenges

ReCAN’s performance on RetinaMNIST (62.3% accuracy on 1,080 training samples) highlights fundamental challenges associated with limited training data. While this performance represents the best reported result for this dataset, the absolute accuracy level remains insufficient for clinical deployment. The limited performance reflects both dataset size constraints and inherent task difficulty—distinguishing between five progressive diabetic retinopathy severity grades requires identifying subtle differences in lesion characteristics that may fall within normal variation.

This observation generalizes beyond RetinaMNIST to the broader challenge of rare disease classification. Many serious medical conditions affect only small patient populations, making it difficult to assemble training datasets of sufficient scale. Standard supervised learning paradigms struggle in these data-scarce regimes, suggesting that alternative approaches may prove necessary. Several promising directions merit investigation:

Few-shot learning frameworks train models to learn from limited examples by explicitly optimizing for rapid adaptation to new tasks. Meta-learning approaches such as MAML or Prototypical Networks could enable ReCAN to leverage knowledge from common diseases (where large datasets exist) to improve rare disease classification with minimal disease-specific training data. Incorporating few-shot learning capabilities would require architectural modifications—perhaps adding task-specific attention modules that rapidly adapt to new diagnostic categories—but could dramatically improve data efficiency.

Semi-supervised and self-supervised learning leverage large quantities of unlabeled medical images to learn robust feature representations without requiring expensive expert annotations. Recent advances in contrastive learning, masked autoencoding, and consistency regularization have demonstrated impressive results on natural images. Adapting these techniques to medical imaging—perhaps through domain-specific augmentation strategies or pretraining objectives aligned with medical image characteristics—could reduce dependence on labeled data while improving feature quality.

Active learning strategies optimize data annotation budgets by strategically selecting which examples to label based on expected information gain. Rather than randomly sampling images for annotation, active learning identifies examples where the model remains most uncertain, maximizing label efficiency. Integrating uncertainty quantification into ReCAN through techniques such as Monte Carlo dropout or deep ensembles would enable active learning deployment, potentially reducing annotation costs while improving performance on underrepresented categories.

5.10.3 Extension to 3D Medical Imaging

Modern medical imaging increasingly involves volumetric modalities—CT, MRI, and 3D ultrasound—that provide rich spatial context unavailable in 2D slices. ReCAN’s

current 2D architecture cannot directly leverage this volumetric information, limiting applicability to important clinical domains. Extending to 3D presents both opportunities and challenges that require careful consideration.

The most straightforward extension involves replacing 2D convolutions with 3D convolutions throughout the architecture, creating a 3D ResNet backbone augmented with 3D channel attention modules. However, this approach increases computational cost and memory requirements by approximately an order of magnitude—3D convolutions operate on volumetric neighborhoods rather than planar neighborhoods, dramatically increasing FLOPs. The resulting model may prove too expensive for deployment scenarios where ReCAN’s efficiency currently provides advantages.

Alternative approaches balance 3D modeling capabilities with computational constraints through various architectural strategies

Chapter 6

Conclusion

6.1 Summary of Key Findings and Implications

This thesis presented ReCAN (Residual Channel Attention Network), an efficient convolutional neural network architecture for medical image classification. Our model achieved 95.47% accuracy on the PathMNIST dataset while maintaining computational efficiency with only 12.9M parameters and 1.7 GFLOPs. Through systematic experimentation with data augmentation, feature enhancement, attention mechanisms, and skip connections, we demonstrated that strategically designed CNNs can match or exceed the performance of more complex hybrid architectures. The key innovation of ReCAN lies in its multi-stage channel attention mechanism with alternate skip connections, which enables effective feature recalibration without the computational overhead of spatial attention or transformer-based approaches. Our results validate that domain-appropriate architectural simplicity, combined with proper regularization and transfer learning, can achieve state-of-the-art performance on medical imaging tasks.

The literature review reveals several critical insights that informed ReCAN’s design:

- **Finding 1: CNNs Remain Competitive:** Despite transformer and state-space model advances, CNNs with attention mechanisms match or exceed their performance on medical image classification when data is limited (<1M images), as is typical in medical domains. This validates ReCAN’s CNN-centric approach.
- **Finding 2: Channel Attention Provides Optimal Cost-Benefit:** Among attention mechanisms (spatial, channel, self-attention), channel attention offers the best accuracy improvement per parameter/FLOP added, particularly for medical images where channel diversity is high and spatial attention’s benefits are marginal for pre-cropped, centered images.
- **Finding 3: Transfer Learning is Essential:** ImageNet pretraining consistently outperforms random initialization for medical image classification, with benefits persisting even on large medical datasets (100K+ images). ReCAN leverages pretrained ResNet-18, aligning with transfer learning best practices.
- **Finding 4: Regularization Trumps Architecture:** For medical imaging, aggressive regularization (heavy dropout, BatchNorm, early stopping) often

improves generalization more than architectural innovations. ReCAN’s MLP design reflects this domain-specific priority.

- **Finding 5: Benchmarking Requires Standardization:** Inconsistent data augmentation, preprocessing, and evaluation protocols hinder fair model comparison. ReCAN adopts MedMNIST’s standardized protocol and MedMamba’s augmentation-free approach to enable direct comparison.
- **Finding 6: Complexity Does Not Guarantee Superiority:** Hybrid architectures (MedMamba’s state-space models, ConvNeXt’s modernized design) introduce implementation complexity without consistently outperforming simpler attention-augmented CNNs. ReCAN prioritizes architectural simplicity and reproducibility.

These findings collectively motivate ReCAN as a systematic exploration of efficient CNN design for medical imaging, filling identified research gaps through multi-stage channel attention, strategic skip connections, comprehensive multi-dataset validation, and fair benchmarking practices. The following chapters detail ReCAN’s methodology, experimental results, and implications for medical AI research and clinical deployment.

Performance: ReCAN achieved superior accuracy compared to baseline ResNet-18 (95.64% vs. 93.82%) while maintaining computational efficiency. The model demonstrated strong generalization across various medical imaging modalities in the MedMNIST benchmark.

Efficiency: With 11.4M parameters and 1.8 GFLOPs, ReCAN offers a practical solution for resource-constrained medical settings. The model’s computational requirements are significantly lower than transformer-based alternatives while delivering comparable accuracy.

Design Insights: Channel-only attention proved sufficient for medical image classification where texture and color patterns are diagnostic. Alternate skip connections (applying attention to every other ResNet block) provided optimal accuracy-efficiency balance. Aggressive data augmentation and regularization were essential for preventing overfitting on limited medical datasets.

6.2 Limitations

Despite promising results, several limitations should be acknowledged:

Input Resolution: The model expects 224×224 inputs, requiring downsampling of high-resolution images which may lose fine-grained diagnostic details critical for pathology applications.

Single Modality: ReCAN processes only visual information, whereas clinical diagnosis integrates imaging with patient history, lab results, and other data sources.

Domain Adaptation: While transfer learning from ImageNet was effective, significant domain gaps remain between natural images and specialized medical imaging modalities like electron microscopy or ultrasound.

Clinical Validation: The model has been evaluated only on benchmark datasets and requires extensive validation through prospective clinical trials before deployment in real healthcare settings.

Interpretability: Although channel attention provides some feature-level insights, the model lacks comprehensive explainability mechanisms required for clinical trust and regulatory approval.

6.3 Future Work

Several promising directions for future research include:

Architectural Extensions: Developing 3D versions for volumetric medical imaging (CT/MRI), exploring hybrid channel-spatial attention for tasks requiring precise localization, and implementing adaptive attention mechanisms with learnable reduction ratios.

Multi-Modal Integration: Extending ReCAN to incorporate patient history, lab results, and genomic data alongside imaging for holistic clinical assessment.

Robustness and Generalization: Conducting cross-hospital validation studies, improving adversarial robustness, and developing out-of-distribution detection mechanisms for safer clinical deployment.

Interpretability: Integrating multiple explanation methods (Grad-CAM, SHAP, counterfactual explanations) and validating attention visualizations against expert pathologist annotations.

Clinical Translation: Conducting prospective clinical trials, integrating with hospital PACS and EHR systems, and performing cost-effectiveness analyses to facilitate real-world adoption.

Bibliography

- [1] Y. Yue and Z. Li. “MedMamba: Vision Mamba for Medical Image Classification”. In: *arXiv* (2024). URL: <https://arxiv.org/abs/2403.03849>.
- [2] A. Bhujel et al. “A Lightweight Attention-Based Convolutional Neural Networks for Tomato Leaf Disease Classification”. In: *Agriculture* 12.2 (2022), p. 228. DOI: 10.3390/agriculture12020228.
- [3] S. Pang, Z. Chen, and F. Yin. “Lightweight multi-scale aggregated residual attention networks for image super-resolution”. In: *Multimedia Tools and Applications* 81.4 (2021), pp. 4797–4819. DOI: 10.1007/s11042-021-11138-x.
- [4] Q. Yi et al. “ELANET: Effective Lightweight Attention-Guided network for Real-Time semantic segmentation”. In: *Neural Processing Letters* 55.5 (2023), pp. 6425–6442. DOI: 10.1007/s11063-023-11145-z.
- [5] X. Huo et al. “HiFuse: Hierarchical Multi-Scale Feature Fusion Network for Medical Image Classification”. In: *arXiv* (2022). URL: <https://arxiv.org/abs/2209.10218>.
- [6] Z. Fu, J. Li, and Z. Hua. “MSA-Net: Multiscale spatial attention network for medical image segmentation”. In: *Alexandria Engineering Journal* 70 (2023), pp. 453–473. DOI: 10.1016/j.aej.2023.02.039.
- [7] M. S. Alam et al. “A Multi-Scale context aware attention model for medical image segmentation”. In: *IEEE Journal of Biomedical and Health Informatics* 27.8 (2022), pp. 3731–3739. DOI: 10.1109/jbhi.2022.3227540.
- [8] X. Zhang et al. “Automatic detection of microaneurysms in fundus images based on multiple preprocessing fusion to extract features”. In: *Biomedical Signal Processing and Control* 85 (2023), p. 104879. DOI: 10.1016/j.bspc.2023.104879.
- [9] F. Wang et al. “Residual Attention Network for Image Classification”. In: *arXiv* (2017). URL: <https://arxiv.org/abs/1704.06904>.
- [10] S. Li and J. Huang. “ResGDANet: An Efficient Residual Group Attention Neural Network for Medical Image Classification”. In: *Applied Sciences* 15.5 (2024), p. 2693. DOI: 10.3390/app15052693.
- [11] Y. Zhang et al. “Image Super-Resolution Using Very Deep Residual Channel Attention Networks”. In: *arXiv* (2018). URL: <https://arxiv.org/abs/1807.02758>.
- [12] J. Shao et al. “A Lightweight Convolutional Neural Network Based on Visual Attention for SAR Image Target Classification”. In: *Sensors* 18.9 (2018), p. 3039. DOI: 10.3390/s18093039.

- [13] L. Chen et al. “SCA-CNN: Spatial and Channel-wise Attention in Convolutional Networks for Image Captioning”. In: *arXiv* (2016). URL: <https://arxiv.org/abs/1611.05594>.
- [14] X. Tong et al. “ASCU-Net: Attention Gate, Spatial and Channel Attention U-Net for Skin Lesion Segmentation”. In: *Diagnostics* 11.3 (2021), p. 501. DOI: 10.3390/diagnostics11030501.
- [15] A. Naveed et al. “AD-Net: Attention-based dilated convolutional residual network with guided decoder for robust skin lesion segmentation”. In: *arXiv* (2024). URL: <https://arxiv.org/abs/2409.05420>.
- [16] S. Kumar Sahoo et al. “Self-adaptive moth flame optimizer combined with crossover operator and Fibonacci search strategy for COVID-19 CT image segmentation”. In: *Expert Systems With Applications* 227 (2023), p. 120367. DOI: 10.1016/j.eswa.2023.120367.
- [17] I. D. Mienye et al. “Deep Convolutional Neural Networks in Medical Image Analysis: A Review”. In: *Information* 16.3 (2025), p. 195. DOI: 10.3390/info16030195.
- [18] A. Lou et al. “CaraNet: Context Axial Reverse Attention Network for Segmentation of Small Medical Objects”. In: *arXiv* (2021). URL: <https://arxiv.org/abs/2108.07368>.
- [19] W. Wu et al. “DCANet: Dual Convolutional Neural Network with Attention for Image Blind Denoising”. In: *arXiv* (2023). URL: <https://arxiv.org/abs/2304.01498>.
- [20] S. Ren et al. “Faster R-CNN: Towards Real-Time Object Detection with Region Proposal Networks”. In: *arXiv* (2015). URL: <https://arxiv.org/abs/1506.01497>.
- [21] F. H. Bhoyan et al. “An efficient dual-attention guided deep learning model with interpretability for identifying medicinal plants”. In: *Current Plant Biology* 44 (2025), p. 100533. DOI: 10.1016/j.cpb.2025.100533.
- [22] H. Tang, S. Bai, and N. Sebe. “Dual Attention GANs for Semantic Image Synthesis”. In: *arXiv* (2020). URL: <https://arxiv.org/abs/2008.13024>.
- [23] R. Mahini et al. “Ensemble deep clustering analysis for time window determination of event-related potentials”. In: *Biomedical Signal Processing and Control* 86 (2023), p. 105202. DOI: 10.1016/j.bspc.2023.105202.
- [24] W. Jing et al. “The clinical effectiveness of staple line reinforcement with different matrix used in surgery”. In: *Frontiers in Bioengineering and Biotechnology* 11 (2023), p. 1178619. DOI: 10.3389/fbioe.2023.1178619.
- [25] A. Ahmadi, E. Mohammadnejadi, and N. Razzaghi-Asl. “Gefitinib derivatives and drug-resistance: A perspective from molecular dynamics simulations”. In: *Computers in Biology and Medicine* 163 (2023), p. 107204. DOI: 10.1016/j.compbiomed.2023.107204.