

ENHANCING AUTOMATED MELANOMA DETECTION: COMPARATIVE
ANALYSIS OF CNN ARCHITECTURES USING THE SIIM-ISIC 2020 KAGGLE
DATASET

by

Bhanu Prasad Sheri, B.E. (Mechanical) M.B.A. (MIS)

DISSERTATION

Presented to the Swiss School of Business and Management Geneva

In Partial Fulfillment

Of the Requirements

For the Degree

DOCTOR OF BUSINESS ADMINISTRATION

SWISS SCHOOL OF BUSINESS AND MANAGEMENT GENEVA

DECEMBER, 2025

ENHANCING AUTOMATED MELANOMA DETECTION: COMPARATIVE
ANALYSIS OF CNN ARCHITECTURES USING THE SIIM-ISIC 2020 KAGGLE
DATASET

by

Bhanu Prasad Sheri

Dr. Anna Provodnikova

<Mentor's full official name>

APPROVED BY



Dissertation chair

RECEIVED/APPROVED BY:



Admissions Director

Dedication

This dissertation is wholeheartedly dedicated to the pillars of strength in my life—my beloved parents and mentors—whose unwavering support and belief in my potential made this journey possible.

To my mother and father, your sacrifices, values, and encouragement shaped the foundation of my education and inspired me to pursue excellence. Your love has been my greatest source of strength.

To Dr. Anna Provodnikova, my esteemed advisor, thank you for your visionary guidance and consistent mentorship throughout this research. Your insights were instrumental in helping me navigate the complexities of deep learning and medical imaging.

To all healthcare professionals and patients fighting against melanoma—this work is driven by the hope that technology can bring timely detection and save lives. May this research contribute in some way to easing the burden of this challenging disease.

This accomplishment is not mine alone; it is the result of a collective commitment to knowledge, care, and perseverance.

Acknowledgements

First and foremost, I am profoundly grateful to Dr. Anna Provodnikova, my dissertation chair, for her exceptional mentorship, academic insight, and unwavering encouragement throughout this research journey. Her constructive feedback and commitment to scholarly excellence helped shape this work into its final form.

I extend my sincere thanks to the Swiss School of Business and Management Geneva for providing a rigorous academic environment and the resources necessary to undertake this dissertation. The faculty and staff have been supportive at every stage of the Doctor of Business Administration program.

I am also indebted to the data science and medical imaging communities, especially the organizers of the SIIM-ISIC 2020 Melanoma Classification Challenge, for making large-scale annotated datasets publicly available. Their contributions laid the groundwork for many advancements in AI-based medical diagnostics, including this study.

To my friends and fellow DBA scholars, thank you for the camaraderie, inspiration, and late-night discussions that fuelled my motivation during this journey.

Finally, I wish to express my heartfelt gratitude to my family—your patience, love, and belief in my abilities gave me the strength to persevere. This achievement would not have been possible without your constant support and sacrifice.

ABSTRACT

ENHANCING AUTOMATED MELANOMA DETECTION: COMPARATIVE ANALYSIS OF CNN ARCHITECTURES USING THE SIIM-ISIC 2020 KAGGLE DATASET

Bhanu Prasad Sheri
2025

Dissertation Chair: Dr. Anna Provodnikova

This investigation explores the combination of AI and behavioural analytics for melanoma detection, integrating deep learning model assessment with the statistical analysis of factors influencing user adoption. The work compares the diagnostic efficiency and interpretability of three state-of-the-art convolutional neural networks—EfficientNetB0, ResNet50, and MobileNetV2—using the SIIM-ISIC 2020 dataset. The preprocessing of data was performed by image normalisation, resizing, and augmentation of the images to prevent overfitting and improve generalisation. Transfer learning was done using ImageNet weights, with models trained using the Adam optimizer and categorical cross-entropy loss. The performance of the models was measured in terms of accuracy, precision, recall, F1-score, and AUC-ROC, whereas Grad-CAM visualisation was used to provide information about the localisation of the lesion and understanding. The computational

results were supported by a quantitative study performed using SPSS, which focused on the healthcare professionals' perceptions, attitudes, and willingness to adopt AI-based melanoma detection systems. The survey data were processed by descriptive statistics of 100 respondents, correlation analysis, and hypothesis testing through Spearman's Rank Correlation and regression models. These tests explored the impact of performance expectancy, effort expectancy, social influence, facilitating conditions, and ethical trust on user acceptance. Findings disclosed significant positive correlations between perceived usefulness, trust, and adoption intention that were consistent with the human-centric dimension of AI integration in the clinical context. In general, the EfficientNetB0 model was able to obtain the top AUC-ROC score of 0.5523; however, all of the models showed low sensitivity towards melanoma because of the imbalance of the dataset. The results emphasise that, on the technical side, performance alone is not enough for successful implementation; proper training, sophisticated loss functions (like focal loss), and ethical transparency are still necessary. The deep learning and SPSS studies, taken together, suggest that the dermatology AI intervention will be effective only if the algorithmic issues are resolved and there is social, ethical, and professional agreement with the clinical stakeholders.

TABLE OF CONTENTS

List of Tables	ix
List of Figures	xi
List of Abbreviations	xii
CHAPTER I: INTRODUCTION.....	1
1.1 Background and Drive	1
1.2 Research Objectives.....	6
1.3 Purpose of Research.....	12
1.4 Significance of the Study	14
CHAPTER II: REVIEW OF LITERATURE	17
2.1 Evolution of Deep Learning in Medical Imaging.....	17
2.2 CNN Architectural Developments	21
2.3 Current Research Limitations	27
CHAPTER III: METHODOLOGY	38
3.1 Research Design Framework	38
3.2 Secondary Data Collection Procedure	44
3.3 Data Preprocessing and Architecture Implementation	51
3.4 Primary Data Collection Procedure	57
3.5 Performance Evaluation Framework for Secondary Data Analysis.....	59
3.6 Primary Data Analysis Using Statistical Package for Social Science (SPSS) and Smart-PLS.....	65
CHAPTER IV: RESULTS.....	68
4.1 Overall Performance Summary.....	68
4.2 Visual Analysis Overview	74
4.3 Detailed Performance Analysis by Architecture.....	78
4.4 Class Imbalance Impact Analysis	87
4.5 Statistical Significance Analysis.....	91
4.6 Survey Analysis Findings and Model Visualisation.....	99
CHAPTER V: DISCUSSION.....	118
5.1 Interpretation of Results.....	118
5.2 Comparison with Previous Studies	125
5.3 Limitations and Challenges.....	129
5.4 Clinical and Practical Implications	134

CHAPTER VI: CONCLUSIONS AND RECOMMENDATIONS.....	140
6.1 Summary of Key Findings	140
6.2 Implications for Medical AI Development	144
6.3 Specific Recommendations.....	148
6.4 Future Research Directions.....	152
6.5 Final Conclusions.....	154
REFERENCES	158
APPENDIX A:.....	181

LIST OF TABLES

Table 1.1: Overview of Research Purpose.....	12
Table 1.2: Significance of the Study.....	15
Table 3.1: Survey Structure and Key Constructs.....	58
Table 4.1: Comprehensive Model Performance Summary.....	68
Table 4.2: Bootstrap Confidence Intervals for Primary Metrics.....	69
Table 4.3: Paired Comparison Test Results.....	70
Table 4.4: Multiple Comparison Correction Results	71
Table 4.5: EfficientNetB0 Detailed Performance Analysis.....	78
Table 4.6: ResNet50 Detailed Performance Analysis.....	81
Table 4.7: MobileNetV2 Detailed Performance Analysis	83
Table 4.8: Bootstrap Confidence Intervals for Primary Metrics.....	91
Table 4.9: Paired Comparison Test Results.....	92
Table 4.10: Multiple Comparison Correction Results	93
Table 4.11: Statistical Power Analysis Results.....	94
Table 4.12: Sensitivity Analysis Results	95
Table 4.13: Clinical Significance Assessment.....	97
Table 4.14: Reliability Statistics.....	99
Table 4.15: Demographic Characteristics of Respondents	99
Table 4.16: Perceptions Toward AI-Based Melanoma Detection	102
Table 4.17: Perceptions Toward Effort Expectancy of AI-Based Melanoma Detection Systems.....	104
Table 4.18: Perceptions Toward Social Influence on the Adoption of AI-Based Melanoma Detection Systems.....	105
Table 4.19: Perceptions Toward Facilitating Conditions for AI-Based Melanoma Detection Systems.....	106
Table 4.20: Perceptions Toward Trust and Reliability of AI-Based Melanoma Detection Systems.....	108
Table 4.21: Perceptions Toward Ethical and Privacy Concerns in AI-Based Melanoma Detection Systems.....	109
Table 4.22: Perceptions Toward Behavioral Intention to Use AI-Based Melanoma Detection Systems.....	111

Table 4.23: Descriptive Statistics	113
Table 4.24: Spearman’s Rank Correlation among Study Variables	114
Table 4.25: Spearman’s Rank Correlation among Social Influence, Facilitating Conditions, and Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies	115
Table 4.26: Correlation between Trust in AI Systems, Ethical and Privacy Concerns, and Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies	116

LIST OF FIGURES

Figure 3.1: Conceptual Model	67
Figure 4.1: Comprehensive Model Performance	68
Figure 4.2: EfficientNetB0, ResNet50 and MobileNetV2 Grad-CAM Visualization	74
Figure 4.3: (a) EfficientNetB0 Confusion Matrix, (b) EfficientNetB0 ROC Curve	78
Figure 4.4: (a) ResNet50 Confusion Matrix, (b) ResNet50 ROC Curve.....	80
Figure 4.5: (a) MobileNetV2 Confusion Matrix, (b) MobileNetV2 ROC Curve.....	83
Figure 4.6: Demographic details.....	101
Figure 4.7: Perceptions Toward AI-Based Melanoma Detection.....	103
Figure 4.8: Perceptions Toward Effort Expectancy of AI-Based Melanoma Detection Systems.....	104
Figure 4.9: Perceptions Toward Social Influence on the Adoption of AI-Based Melanoma Detection Systems.....	106
Figure 4.10: Perceptions Toward Facilitating Conditions for AI-Based Melanoma Detection Systems.....	107
Figure 4.11: Perceptions Toward Trust and Reliability of AI-Based Melanoma Detection Systems.....	109
Figure 4.12: Perceptions Toward Ethical and Privacy Concerns in AI-Based Melanoma Detection Systems.....	110
Figure 4.13: Perceptions Toward Behavioral Intention to Use AI-Based Melanoma Detection Systems.....	112
Figure 4.14: PLS-SEM Model	117

LIST OF ABBREVIATIONS

Abbreviations	Full Form
AUC	Area Under the Curve
CNNs	Convolutional Neural Networks
Grad-CAM	Gradient-Weighted Class Activation Mapping
GPU	Graphics Processing Unit
AI	Artificial Intelligence
MCC	Matthews Correlation Coefficient
FDR	False Discovery Rate
SPSS	Statistical Package for The Social Sciences
PLS-SEM	Partial Least Squares Structural Equation Modeling
NND	Number Needed to Diagnose
TAM	Technology Acceptance Model
UTAUT	Unified Theory of Acceptance and Use of Technology

CHAPTER I: INTRODUCTION

1.1 Background and Drive

1.1.1 The Global Melanoma Challenge

Melanoma is another skin cancer type, and it is one of the most aggressive ones: on the one hand, melanoma accounts only for about 4 percent of all skin cancer cases, and, on the other hand, melanoma leads to about 75 percent of skin cancer deaths across the globe (Esteva *et al.*, 2017; Siegel *et al.*, 2023a). This harsh comparative feature highlights the utmost importance of early diagnosis, where the positive five-year outcomes are above 99 and close to 27 percent, respectively, in early-stage and late-stage cases of melanoma (Balch *et al.*, 2009; Society, 2023).

Melanoma levels have been gravely developing across the globe, and the rate of new cases is growing quicker than all other forms of cancer during the last 30 years (Erdmann *et al.*, 2013; Garbe *et al.*, 2016). The life time risk of developing melanoma has been magnified twenty-fold, placing it at the rate of about 1 in 38 as compared to 1 in 600 back in 1960 and therefore this cannot be explained by enhanced reporting or inflated screening procedures (Siegel *et al.*, 2023a; Rastrelli *et al.*, 2014).

Such an epidemiological pattern is evidence of a combination of factors such as predisposition, exposure to ultraviolet radiations, and changing lifestyle designs that have so far established a burning public health necessity that must be addressed in a timely manner (Dennis *et al.*, 2008; Whiteman, Green and Olsen, 2016). That the cost of melanoma to the economy goes much further than those associated with its treatment proper, encompassing long-term care, reduced productivity, as well as the exorbitant toll of premature deaths (Guy *et al.*, 2015).

The National Cancer Institute states that the aggregate cost of treating melanoma in the United States amounts to over 8.1 billion, and the weight of this treatment differs considerably regarding disease progression: 1,732 per patient in case of an early disease, and much more than 159,808 in cases of late-stage metastatic melanoma (Guy *et al.*, 2015). These numbers do not just represent the monetary consequence but also the possible usefulness of an improvement in early detection technologies in moving the diagnosis to a more treatable stage.

1.1.2 Current Diagnostic Paradigm and Limitations

The conventional methods of diagnosis are based on the ABCDE rule (Asymmetry, Border irregularity, Colour variation, Diameter >6mm, Evolving features) and skill of clinical information of dermatologists (Thomas *et al.*, 1998; Argenziano *et al.*, 2003). Nonetheless, this paradigm of human-centeredness has its strong restrictions which limit its applicability and reach in various medical contexts.

The critically supplied shortage of dermatologists in the world is the most important issue. With over 330 million individuals and 12,000 practicing dermatologists, United States has an average of 35 days of wait time to get an appointment with a dermatology physician (Kimball and Resneck, 2008). In the course of this wait time, aggressive melanoma might advance considerably. The reality is even worse in underprivileged communities and underrepresented regions since certain areas do not offer any practicing dermatologists within 100 miles, and hence, the provision of specialist services is excruciatingly hard (Feng *et al.*, 2018; Glazer *et al.*, 2017).

Another major limitation is the possibility of inter-observer variability, as studies indicate that it will be with a diagnostic accuracy below 80 percent efficiency even among skilled dermatologists in a primary clinical assessment (Haenssle *et al.*, 2018; Titus J. Brinker *et al.*, 2019). This has been due to small and unnoticeable morphological

differences between benign pigmented lesions and early melanoma, impact of various demographic factors of patients on the appearance of a lesion, and the subjective nature of visual pattern recognition (Carli *et al.*, 2004; Menzies *et al.*, 1996).

It has been shown that, even veteran dermatologists can give varying diagnoses on one and the same lesion, and how the need of a more objective assessment matter is manifested (Elmore *et al.*, 2017; Marchetti *et al.*, 2018). Due to the growing awareness of unusual presentations and non-classic subtypes of melanoma that do not fit classic ABCDE criteria, the problem of diagnosis is complicated.

Amelanotic melanoma occurs in 5-15 percent of all cases and often cannot be diagnosed properly because of the similarity to benign inflammations (Pizzichetta *et al.*, 2004; Jaimes *et al.*, 2012). Nodular melanomas can also be characterised by a rapid growth where there is no normal radial growth phase so these are difficult to be detected by using standard surveillance methods (Chamberlain *et al.*, 2003).

1.1.3 The Promise of Artificial Intelligence

The use of artificial intelligence technologies, especially when coupled with large and high-quality medical datasets of images, can bring unprecedented capabilities to deal with these problems of dermatological diagnosis and access to dermatological care as much as it is currently available to experts (LeCun, Bengio and Hinton, 2015; Rajpara *et al.*, 2009). The convolutional neural networks have proven to be very effective in medical images processing, with some publications finding that AI systems were as effective or more so than board-certified dermatologists under controlled conditions.

Esteva *et al.* (2017) landmark study held a turning point in AI dermatological research as it showed that networks based on CNNs that had been trained in more than 129,000 clinical images could classify skin lesions with equal or higher accuracy than such board-certified dermatologists. The study introduced the concept not only that AI-based

skin cancer diagnosis is technically feasible but also accelerated the development of systems concerning other parts of the process such as lesion segmentation, risk stratification, and treatment planning (Codella *et al.*, 2018; Tschandl *et al.*, 2020).

Further research by others has confirmed these results using a wider variety of populations, imaging techniques and clinical contexts. A large number of dermatologists was used in an in-depth comparison made by Haenssle *et al.* (2018) and the CNN performed better at recognising melanoma, with the area under the curve (AUC) of 0.86 as opposed to 0.79 in dermatologists. These findings were especially quite impressive since dermatologists at different levels of experience, including residents and department heads, were used, implying that AI systems may always give the same performance, independent of the knowledge of the person operating them.

In addition to the increase in accuracy, such AI-guided diagnosis might create accessibility, consistency, and scalability in dermatological care, which can even alter the paradigm of care provision (Yu *et al.*, 2018); Thomsen *et al.*, 2020). AI systems may offer continuous coverage (24/7), as well as provide the benefits of telemedicine (no geographical limits to care) and standardised assessment protocols, which minimise variability between providers.

AI tools have the ability to recognise patterns even imperceptible to the human eye, to detect an onset of morphological changes that are not visible yet, and combine the data provided by a wide scope of sources such as the patient history, a genetic predisposition to a disease, and environmental exposure (Titus J Brinker *et al.*, 2019); Phillips *et al.*, 2019). The possibility of AI to complement and not to substitute human experience is a paradigm game-changer in the perspective of integrative human-AI diagnostics.

1.1.4 Research Gap and Motivation

Notwithstanding the individual studies with promising outcomes, and despite appropriate methods of study selected, methodological heterogeneity is a major ill in present scope of AI investigations in detection of melanoma (Winkler *et al.*, 2019). Available studies use different datasets, performance factors, experiment conditions, and evaluation procedures, and architectural comparisons are almost impossible, and they make evidence-based decisions toward clinical usage challenging.

This division has established a number of serious gaps in knowledge and this dissertation is trying to fill it in. First, no universal evaluation models exist that would be able to consistently test various CNN components in the same environment (Topol, 2019; Sendak *et al.*, 2020). Second, the majority of research studies attributed all their results to the accuracy measures without attributing, or even evaluating, the deployment issues in terms of computation usage, interpretability features, etc (Rajkomar *et al.*, 2018; Liu *et al.*, 2020).

Third the existing studies do not sufficiently embrace the impact of imbalanced classes to sufficient degree and such a factor is a ubiquitous issue in medical data where the number of pathological cases is significantly lower than the normal ones (Krawczyk, 2016; Johnson *et al.*, 2019). The high degree of class imbalance that is common in melanoma screening (a 50:1 ratio of benign lesions to melanoma or higher), poses essentially intractable difficulties to using standard machine-learning methods which are often not dealt with very methodically.

The realisation that the effective use of AI to detect melanoma cannot be accomplished solely through technical innovation, but rather through rigorous comparative assessment that can be used to guide the development of an evidence-based practice, drives the need to conduct this study (Char *et al.*, 2018)(Finlayson *et al.*, 2021). The following is

a clear, scientifically backed set of advice to providers, technology developers, and regulators of which AI solutions are best suited to various clinical settings and how to assess them based on their accuracy requirement, the limitations of their demands on computation, the interpretability demands and the setting of implementation.

Moreover, the escalating situation with the epidemic of melanoma, on the one hand, and the ongoing inability of conventional diagnostic methods to address the problem, as well as the potential and proven efficacy of AI technologies, on the other hand, presents a strong argument in favour of faster but properly conducted research that could fill the gap between laboratory and clinical practice (Beam and Kohane, 2018; Obermeyer and Emanuel, 2016). As a dissertation, the work reflects a deliberate attempt to supply the comparative analysis and implementation advice required to develop the discipline into feasible, useful to the clinic AI-assisted melanoma detection systems.

1.2 Research Objectives

1.2.1 Primary Research Objectives

This dissertation can fill the most crucial gaps in the field of research on AI-assisted diagnosis methods of melanoma by comparing the best CNN systems developed in an automated and controlled experiment. These main aims are meant to create practical knowledge that can be used in practice regarding the choice of medical AI implementation, as well as to set methodological benchmarks with the examples to be used in any further scientific comparative studies at this overlapping field.

- **Objective 1: Comprehensive Architectural Comparison**

The first objective aims at running a controlled and consistent comparison of EfficientNetB0, ResNet50 and MobileNetV2 under the same preprocessing, experimental and evaluation conditions to overcome methodological inconsistency in the current literature, which complicates reliable building comparisons. It is possible to measure

important performance metrics, including accuracy, sensitivity, specificity, computational efficiency, and clinical relevance with the help of rigorous statistical methods, such as bootstrap confidence intervals, cross-validation, and significance testing to determine whether there is any meaningful difference (Rajpurkar *et al.*, 2017; McKinney *et al.*, 2020). This methodology allows ranking architectural features based on evidence and facilitates the process of making informed choices in deployments and designing future research, which is no longer at anecdotal or inconsistent (Nagendran *et al.*, 2020; Liu *et al.*, 2019).

- **Objective 2: Clinical Interpretability Assessment**

The scope of Objective 2 is to determine a systematic assessment of the interpretability of any given architecture based on the visualization quality and clinical relevance, considering that AI in healthcare should be evaluated by performance as well as explainability and reliability to the medical practitioner (Holzinger *et al.*, 2017; Adadi and Berrada, 2018). Learned attention patterns will be visualised via Grad-CAM, and the clinical meaningfulness of the extracted features will be evaluated, as well as to verify whether the model depends on good cues or coincidences, across the types of lesions (Selvaraju *et al.*, 2017; Chattopadhyay *et al.*, 2018). It is a framework of interpretability that is supposed to help develop clinically actionable modes of explanation that complement, rather than substitute, clinical decision-making in the larger context of explainable medical AI (Tjoa and Guan, 2021; Murdoch *et al.*, 2019).

- **Objective 3: Computational Resource Analysis**

Objective 3 looks at the computational values of each architecture: training time, the inference latency, memory usage, and energy consumption, to determine the effect of resource requirements on feasibility, cost-effectiveness, and scalability in diverse clinical settings (Strubell *et al.*, 2019; Schwartz *et al.*, 2020). As the analysis evaluates performance with different hardware, including high-end GPU servers to mobile devices, it is able to

determine the deployment suitability to execute operations under various operational constraints and which models are most suitable in settings at both ends of the spectrum, such as high-throughput screening facilities and point-of-care mobile apps (Thompson *et al.*, 2023; Patterson *et al.*, 2021).

- **Objective 4: Methodological Framework Development**

The goal of objective 4 is to create a consistent methodological framework capable of fostering consistency and comparability of future medical AI studies that will enhance more consistent science in the field (Mongan *et al.*, 2020; Beam *et al.*, 2020). This involves the establishment of templates and protocols in partitioning the datasets, measures of evaluation, statistical analysis and presentation of the results, which will permit more rigorous studies and support more robust meta-analyses and systematic reviews. This type of framework will add a sustainable methodological value that will enhance the capacity of the research community to produce credible, generalizable research on the efficacy of medical AI.

1.2.2 Secondary Research Objectives

- **Investigation of the Impact of Class Imbalance**

To study the impact of extreme class imbalance in CNN learning behaviour through the analysis of learning curves, convergence, and performance related to trends, architectural features that provide higher resilience to imbalanced data are sought (He and Garcia, 2009; Krawczyk, 2016).

- **Testing the performance of Transfer Learning**

In order to assess the effect of transfer learning between architectures by examining the effect of pre-training on natural or medical-domain images on features transferability, fine-tuning, domain adaptation, and final performance on medical imaging tasks (Raghu *et al.*, 2019; Azizi *et al.*, 2021).

- **Deployment Readiness Clinical Assessment**

In order to determine the suitability and efficiency of the chosen CNNs in a real clinical practice with different quality of images, diverse patient groups, and different equipment, it is necessary to determine the limitations that should be considered prior to clinical implementation (Keane *et al.*, 2018; McKinney *et al.*, 2020).

- **Analysis of Errors and Development of Mitigation Strategy**

To examine modes of architectural failure, that is, false positives and false negatives, to study the patterns of errors and construct training strategies that enhance reliability in clinical practice (Larrazabal *et al.*, 2020; Oakden-Rayner *et al.*, 2020).

1.2.3 Behavioural and Acceptance-Oriented Research Objectives

- **Performance and Effort Expectancy Analysis:** To determine the effects of performance and effort expectancy on the acceptance of the automated melanoma detection system by healthcare professionals and patients.
- **Social Influence and Facilitating Conditions:** To examine the influence of social pressure and the presence of favourable circumstances in the formation of the attitudes of users to the use of such AI diagnostic tools.
- **Trust and Ethical Considerations:** To determine the effects of trust in AI, ethical and privacy issues on acceptance of automated melanoma detection technologies.
- **Demographic Moderation Effects:** To evaluate the moderating effect of demographic factors (such as age, gender, education, and profession) on the relationship between key predictors and the acceptance of automated melanoma detection systems.

1.2.4 Research Questions

Research objectives convert to research questions that can be answered by experimental design, methods of data analysis:

1. **Performance Compared:** Which CNN architecture shows the better performance on melanoma detection when tested in the same conditions as other using clinically relevant metrics?
2. **Interpretability Analysis:** What is the comparative mix of interpretability attributes of various architectures, and which of those methods represent the most clinically useful interpretation of attribution of features?
3. **Resource Efficiency:** What are the trade-offs in terms of computations between the various architectures and what are their effects on the applicability to various healthcare environments in terms of deployment?
4. **Class Imbalance Robustness:** How do these architectures perform under extreme class imbalance, and which architectural features are the most desirable in providing robustness to imbalanced medical data in general?
5. **Clinical Translation:** What do we see as the biggest obstacles to getting existing CNN methods to the clinic and what needs to be changed to pass with clinically acceptable performance?

1.2.5 Success Criteria

The following several criteria that will prove the technical success and practice appeal of the research will be used to assess the success of it:

1. **Technical Success criteria:**
 - Effective use and/or comparison of the three CNN architectures with the same experimental settings

- Granting of statistically significant levels in performance comparisons with confidence levels of at most 5 %
- The finished interpretability analysis, which comprises the quantitative evaluation of the consistency of attentional patterns
- Broad coverage of computation infrastructure on various hardware facilities

2. Clinical Relevance Criteria:

- Determination of architectures showing clinically significant sensitivity (over 80%) to detect melanoma
- Formulation of deployment suggestion according to the performance-resource trade failures
- Creation of interpretability rules that may be used in clinical decision-making
- Preparation of assessment tools that emphasize safety and clinical applicability of patients

3. Criteria of Methodological Rigor:

- Observation of accepted procedures of statistical analysis and multiple comparisons corrections
- Deployment of stringent cross validation measures using external validation data sets
- Record of each practice of the experiment with all the adequacy to reproduce the procedure
- Integration of uncertainty quantification to all performance assessment

Such goals and criteria of success will serve as a thorough guideline in carrying out rigorous and clinically applicable studies that will ultimately contribute to advancing the

field of AI-assisted melanoma detection, as well as set standards in a future comparative analysis in the field of medical AI applications.

1.3 Purpose of Research

The purpose of the study is to meet high-stakes clinical, scientific, methodological, and policy requirements concerning the design and implementation of CNN-based melanoma detection systems. It intends to give evidence-based architectural comparisons, enhance the knowledge of model behaviour with clinical constraints, enhance interpretability requirements, and establish methodological uniformity in further medical AI studies. It is also useful in the practical application of the work, as it analyses resource requirements, clinical feasibility, and stakeholder needs, and it also advances the larger objectives of health equity, preparedness in the regulatory sphere, and informed policymaking. The study provides information that is useful to clinicians, developers, healthcare administrators, regulators, and patients in a manner that forms the future AI systems, which are scientifically sound, clinically relevant, and usable in a wide variety of settings.

Table 1.1: Overview of Research Purpose

Key Purpose	Justification
Addressing Critical Clinical Needs	- Establishes evidence-based standards for CNN selection, reducing confusion caused by conflicting research (Singh <i>et al.</i>, 2021; Wiens <i>et al.</i>, 2019). - Highlights rising melanoma rates and limited dermatology access, reinforcing need for validated AI tools (Glazer <i>et al.</i>, 2017). - Offers cost-benefit insights for healthcare administrators regarding performance and resource demands (Ahuja, 2019; Chen, Liu and Peng, 2019). - Provides interpretability and clinical utility guidance to support practitioner trust (Cabitza <i>et al.</i>, 2017; Chen <i>et al.</i>, 2020).

	• Gives developers benchmarks to align innovation with clinical relevance (Sendak <i>et al.</i>, 2020; (Shah, Milstein and Bagley, 2019).
Advancing Scientific Understanding	• Contributes knowledge on how CNN architectures behave with highly imbalanced medical datasets (Johnson <i>et al.</i>, 2019; Buda, Maki and Mazurowski, 2018). • Examines impact of architectural elements such as residual connections and depthwise separable convolutions (Ker <i>et al.</i>, 2018; Shen, Wu and Suk, 2017). • Studies interpretability to support explainable, clinically verifiable AI systems (Holzinger <i>et al.</i>, 2017; Samek and Müller, 2019). • Lays groundwork for integrated human–AI decision-making workflows (Topol, 2019; Yu <i>et al.</i>, 2018).
Establishing Methodological Standards	• Addresses methodological inconsistency in medical AI by proposing unified protocols (Varoquaux and Cheplygina, 2022; Winkler <i>et al.</i>, 2019). • Provides templates for dataset splitting, preprocessing, model training, evaluation, and statistical analysis. • Supports regulatory and clinical validation requirements through consistent evaluation frameworks (Habli, Lawton and Porter, 2020; Ahmad <i>et al.</i>, 2021).
Facilitating Evidence-Based Implementation	• Evaluates computational requirements, interpretability, workflow integration, and resource usage for real-world deployment (Keane <i>et al.</i>, 2018; Nagendran <i>et al.</i>, 2020). • Supports organisational decision-making by clarifying trade-offs in cost, infrastructure, staff training, and clinical value (Chen, Liu and Peng, 2019; Ahuja, 2019). • Offers practical recommendations to improve technical readiness for clinical use.
Contributing to Health Equity and Access	• Aims to democratise dermatological expertise, addressing disparities in underserved populations (Adamson and Smith, 2018; Rajpara <i>et al.</i>, 2009).

	• Evaluates architectures suitable for low-resource environments, including mobile and telemedicine settings (Udrea <i>et al.</i>, 2020; Finnane <i>et al.</i>, 2017). • Emphasises computational efficiency to enable deployment on affordable hardware (Thompson <i>et al.</i>, 2022; Strubell <i>et al.</i>, 2019).
Supporting Policy and Regulatory Development	• Provides evidence needed for regulatory agencies assessing AI-based medical devices (Reddy <i>et al.</i>, 2020; Matheny, Whicher and Israni, 2020). • Establishes assessment standards for evaluating safety, performance, and clinical relevance (Habli, Lawton and Porter, 2020; Ahmad <i>et al.</i>, 2021). • Informs healthcare policy, insurance decisions, and technology investment through cost–benefit and resource analysis (Matheny, Israni and Ahmed, 2018; Price, Gerke and Cohen, 2019).

1.4 Significance of the Study

This study is important since it contributes in a multifaceted way to the clinical practice, scientific methodology, policy of health care, economic assessment, education and development of technology in the long run. The research will enable policy makers, clinicians, and healthcare organizations make informed decisions about regulation, investment, and implementation of AI-assisted melanoma detection systems by offering evidence-based information on their performance, safety, interpretability, and resource requirements. Its uniform methodological framework improves the quality, comparability and reliability of future medical AI studies, and its studies of class imbalance, computational needs and workflow combination improve the scientific knowledge required to develop clinically relevant and sustainable AI systems. What is more, the work can be used as a helpful educational material by medical workers and scientists and serve as a foundation of future developments in the medical imaging field and human-AI interaction that has a broad and long-term effect on the medical and the technological community.

Table 1.2: Significance of the Study

Key significance	Shortened Content
Clinical Impact and Patient Safety	The analysis offers facts that inform policymakers and regulators on the application of AI in healthcare (Reddy <i>et al.</i> , 2020; Matheny, Whicher and Israni, 2020). The procedures designed here aid in the regulatory documentation by evaluating technical, clinical, and safety-related factors of AI systems (Habli, Lawton and Porter, 2020; Ahmad <i>et al.</i> , 2021). Knowledge about the failure modes and risk of accidents can be used to inform the oversight mechanisms. The studies also justify the investment in healthcare, insurance policy, and professional standards based on the cost-benefit analysis and resource analysis (Matheny, Israni and Ahmed, 2018; Price, Gerke and Cohen, 2019). The framework is useful in scientific advancement, clinical implementation, and policy formulation, which would facilitate responsible integration of AI-assisted melanoma detection.
Methodological Innovation and Scientific Advancement	The study advances medical AI methodology by addressing weaknesses in current evaluation practices (Chen and Asch, 2017; Varoquaux and Cheplygina, 2022). It offers a standardised framework that improves research quality and comparability. Performance, interpretability, and computational demands are assessed under unified conditions to support more meaningful architectural comparisons (Liu <i>et al.</i> , 2019; Nagendran <i>et al.</i> , 2020). The class-imbalance analysis provides insights applicable to rare-disease detection (Buda, Maki and Mazurowski, 2018; Johnson <i>et al.</i> , 2019). Integrating computational resource analysis with clinical evaluation offers a more realistic model of deploy ability (Thompson <i>et al.</i> , 2022; Strubell <i>et al.</i> , 2019).
Economic and Healthcare System Impact	The study supports evidence-based decisions for healthcare technology investment (Ahuja, 2019; Chen, Liu and Peng, 2019). It provides cost-benefit analyses that help organisations optimise clinical value while minimising financial risk. Examining computational requirements informs the total cost of ownership, including hardware, energy use, maintenance, and training (Patterson <i>et al.</i> , 2021; Schwartz <i>et al.</i> , 2020). Improved diagnostic efficiency through AI offers economic gains via early detection and reduced unnecessary testing (McKinney <i>et al.</i> , 2020; Rodríguez-Ruiz

	<i>et al.</i> , 2019). Evidence on clinical efficacy and cost-efficiency also informs insurance reimbursement policies (Matheny, Israni and Ahmed, 2018; Price, Gerke and Cohen, 2019).
Educational and Professional Development Significance	The study contributes to the aspect of professional training by offering evidence-based information about the abilities and shortcomings of AI diagnostic technologies (Topol, 2019; Masters, 2019). Interpretability and error rates analysis can help developers optimize clinician-specific training to use AI safely (Chen <i>et al.</i> , 2020; Cabitza <i>et al.</i> , 2017). The focus on clinical relevance provides advice to AI developers and researchers (Beam <i>et al.</i> , 2020; Rajkomar <i>et al.</i> , 2018). The methodological frameworks developed can be utilized as the teaching resources of graduate students and the novice researchers in medical AI (Varoquaux and Cheplygina, 2022; Winkler <i>et al.</i> , 2019).
Regulatory and Policy Implications	The study offers evidence to inform regulatory evaluation and authorization of AI-based medical devices (Ahmad <i>et al.</i> , 2021; Habli, Lawton and Porter, 2020). Its holistic evaluation strategy can act as a model for regulatory assessment protocols addressing accuracy, safety, interpretability, and deployment feasibility (Reddy <i>et al.</i> , 2020; Matheny, Whicher and Israni, 2020). Findings on the limitations of accuracy-only assessment support the development of medical-specific evaluation tools (Chen and Asch, 2017; Varoquaux and Cheplygina, 2022). The study also supports international regulatory harmonization for AI in melanoma detection (Matheny, Israni and Ahmed, 2018; Price, Gerke and Cohen, 2019)
Long-term Scientific and Technological Impact	The research informs future AI innovation by identifying architectural characteristics best suited for medical imaging (Ker <i>et al.</i> , 2018; Shen, Wu and Suk, 2017). Insights into class imbalance and mitigation strategies support development of AI for detecting rare conditions (Buda, Maki and Mazurowski, 2018; Johnson <i>et al.</i> , 2019). Work on interpretability and workflow integration guides future human-AI collaboration models (Topol, 2019; Yu <i>et al.</i> , 2018). The standardized methods and evaluation frameworks lay the foundation for rigorous future studies and robust meta-analyses, supporting long-term advancements in medical AI.

CHAPTER II:

REVIEW OF LITERATURE

2.1 Evolution of Deep Learning in Medical Imaging

2.1.1 Historical Foundation and Early Development

The introduction of deep learning methods to medicine imaging is one of the most revolutionary trends in contemporary health technologies which can be traced to the breakthroughs in artificial neural network concept in the 1980s and 1990s (McCulloch and Pitts, 1943; Rosenblatt, 1958). But the early methods were not easily scaled to medical issues because of computational constraints and because there were simply no big, high-quality annotated data sets that could be used to train through complex neural network models.

The interaction of a number of technological, and methodological developments constituted a turning point in the application of deep learning in medical imaging in the early 2010s. Incremental Graphics Processing Unit (GPU) computing overcame the disparity in computation between growing neural networks, which immensely decreased the computational threshold in the training of big neural networks (Chetlur *et al.*, 2014; Jouppi *et al.*, 2017). At the same time, the development of architectures in networks, especially the emergence of convolutional neural networks (CNNs), allowed to find effective applications in image analysis, which would be capable of understanding complex visual structures (LeCun *et al.*, 1989; Fukushima, 1980).

With the breakthrough demonstration that CNNs have transformative potential in Krizhevsky, Sutskever and Hinton (2012) through the AlexNet architecture, which produced record-breaking performance on the ImageNet Large Scale Visual Recognition Challenge, interest grew in machine-learning applications in other areas, to include medical imaging where researchers found possible applications in automated analysis of highly

complex visual patterns, previously not possible in large part due to the level of human expertise previously needed to study them (Deng *et al.*, 2009; Russakovsky *et al.*, 2015).

The first medical imaging use cases of deep learning concentrated on the radiology applications due to the presence of sufficient datasets of annotated images and the existing technical framework of digital image processing (Doi, 2007; Giger, 2018). Early works by Greenspan, van Ginneken and Summers (2016), and Litjens *et al.* (2017) that made a comprehensive review of early applications, which showed promising results in a wide variety of applications with lesion detection, organ segmentation and disease classification in many imaging modalities.

During the formative years, some of the essential principles that govern the development of medical imaging AI became quite evident, such as the value of large and diverse training datasets, the necessity of adapting general imaging computer vision methods to specific fields, and the necessity of clinical validation as the solution to technical breakthroughs applied to reality in healthcare (Rajpurkar *et al.*, 2017; De Fauw *et al.*, 2018).

2.1.2 Breakthrough Studies in Dermatological AI

Utilisation of deep learning in dermatological diagnosis makes a particularly important step because of the strong visual component of skin lesions diagnosis, as well as the existence of high-resolution dermoscopic imaging material (Argenziano *et al.*, 2003; Kittler *et al.*, 2002). This historical study by Esteva *et al.* (2017) shook up the field of dermatological AI by showing that CNN can provide skin cancer classification results similar to humans, addressing, especially certified dermatologists.

The Esteva research took advantage of an unprecedented set of clinical images of more than 129,000 clinical images over 2,032 distinct diseases, which was the largest and most extensive dermatological collection that had ever been compiled (Esteva *et al.*, 2017).

The authors used a modified version of the Inception v3 pre-trained on ImageNet and further fine-tuned to dermatological classification tasks. The findings indicated that the CNN performed at the same level as the 21-board certified dermatologists in various diagnostic tasks such as the detection of melanoma and the identification of basal cell carcinoma.

These results were later expanded upon by other researchers with varying architectures, data set constructions and measures of assessment. Haenssle *et al.* (2018) performed a large comparative test across 58 dermatologists in 17 countries that demonstrated the better performance of melanoma recognition using a CNN than with human experts at various levels of experience and expertise. The demonstration in the study that AI performed better than experienced dermatologists caused considerable interest in the clinical sphere and increased the speed at which new research on practical approaches to deployment was carried out.

Tschandl *et al.* (2019) contributed to the community by developing the HAM10000 dataset, which has already been widely used as a benchmark in dermatological AI research (Tschandl, Rosendahl and Kittler, 2018). The data in this collection overcome most of the shortcomings of previous collections by employing consistent imaging procedures, full metadata annotation, and curation/54 Due to the availability of such standardised datasets, research has become more reproducible and meaningful cross-sectional comparisons between research groups have been made.

These breakthrough studies in the dermatological AI field set a number of precedents used in medical AI research, such as the feasibility of the AI system generating expert-level results on multi-modal and complex diagnostic tasks, the necessity of large training datasets with ample diversity to guarantee successful performance of the system,

the centrality of clinical validation by a wide range of expert reviewers (Titus J. Brinker *et al.*, 2019; Tschandl *et al.*, 2020).

2.1.3 Current State and Emerging Trends

Rapid technological development and the accompanying heightened attention to clinical translation issues and the need to adjust medical imaging-related AI to the real world are the hallmark features of modern deep learning state in medical imaging (Topol, 2019; Yu *et al.*, 2018). More recent studies have abandoned the more technical aspect of minimising performance and moved to focus on the multi-faceted facet of practical deployment requirements.

Extended pre-training methods especially those for medical imaging are one of the new directions as they develop foundation models (Zhou *et al.*, 2020; Azizi *et al.*, 2021). At the same time, these directions are expected to exploit the increased accessibility of medical imaging data to develop more generalizable models whose application to downstream diagnostic tasks can be made efficient. The strategies can assist in overcoming the issue of data scarcity that hinders the use of AI in rare diseases and niche applications.

The next important direction is multi-modal data incorporation and the introduction of systems that incorporate imaging and clinical metadata, patient history, laboratory results, genetic data, and other variables into single systems to achieve a more comprehensive diagnostic system (Huang *et al.*, 2020; Rajpurkar *et al.*, 2022). This holistic strategy is due to the fact that optimal clinical decision-making usually involves using several sources of information, other than imaging analysis.

Federated learning and privacy-preserving methods are becoming more popular because of the greater awareness of the privacy and security concerns of healthcare applications on data (Li *et al.*, 2020; Kaissis *et al.*, 2020). Such methodologies allow the creation of collaborative models between several institutions at the same time, and also do

not violate the confidentiality of the patient and do not contradict the norms of privacy protection, including the provisions of HIPAA and GDPR.

The implementation of quality assurance and frequent surveillance of AI systems at the point of use is an emerging practice that clearly indicates that healthcare institutions are just starting to implement strategies of periodical performance assessment of AI systems in the clinical setting (Chen *et al.*, 2021; Finlayson *et al.*, 2021). These are creation of automated quality control mechanisms, drift-watching algorithms and methods of preserving model performance as data distribution and clinical methodologies change with time.

The history of deep learning in medical imaging shows a trend between initial technical studies carrying out a technical proof of concept to fully implemented systems with potential clinical use. This evolution has been fuelled by the core promise of AI to enhance diagnostic precision, spread access to expert-level evaluation, as well as increase efficiency and effectiveness of healthcare delivery across a wide range of clinical contexts.

2.2 CNN Architectural Developments

2.2.1 ResNet50 Architecture: Foundation of Residual Learning

ResNet (Residual Network) architecture proposed by He *et al.* (2016) is one of the most influential inventions in the field of deep learning since it completely shifts the way very deep neural networks are modelled and trained. A residual connection addressed the vanishing gradient issue that had restricted networks to just a handful of layers, and allowed networks with hundreds or thousands of layers to be built, keeping the gradient flow stable throughout training.

Architecture innovation and Design principles

The novelty of ResNet is that its residual blocks make a network learn residual instead of direct mappings using skip connections (Veit, Wilber and Belongie, 2016; Zhang

et al., 2017). As opposed to learning a direct mapping $H(x)$ in input-output of inputs x , ResNet blocks learn the residual function $F(x) = H(x) - x$, so that the network can learn small adjustments to identity mappings. The method mitigates the degradation issue with deeper networks having poorer training error than shallow networks even though in principle the extra layers ought to be able to support identity mappings of the same properties as the shallow networks.

ResNet50 has a 50-layered structure including 49 convolutional layers and a single fully connected layer that is divided into four residual stages by distinct spatial resolutions (He *et al.*, 2016)(Wu, Shen and van den Hengel, 2019). It starts with a 7×7 convolution layer, which is implemented by batch normalization and activation through ReLU, then max pooling and then the stages of residual blocks. In each stage, there are several residual blocks which preserve the spatial dimensions of the stage whilst allowing the possible alteration of the number of channels.

In ResNet50, a bottleneck structure is used to keep the representational capacity of residual blocks, but lower the computation complexity (He *et al.*, 2016; Zagoruyko and Komodakis, 2017). The bottleneck blocks have three convolutional layers: two convolutional layers (1×1 and 3×3) are used to down sample the channels and extract the spatial features respectively and the residual channels are restored with the help of a 1×1 convolution. This construction allows multi-scale feature extraction with less cost per block to compute.

The residual learning can be described mathematically as $y = F(x, \{W_i\}) + x$ with $F(x, \{W_i\})$ extending to residual mapping and W_i layer parameters(He *et al.*, 2016). It is a formulation where when the identity mappings are considered to be optimal, the residual functions can be forced towards zero and hence is easier to optimize than identity mapping itself.

Medical Imaging Performance and Clinical Applications

ResNet50 has proven a spectacular result in many medical imaging applications which make it a well-known choice to be used as a backbone in the development of clinical AI systems (Rajpurkar *et al.*, 2017; Esteva *et al.*, 2017). The ResNet50-based systems have proven successful in dermatological use and posted over 85 percent accuracy in data over the ISIC 2019 datasets, with efficacy ranging across the various skin lesion types and imaging conditions (Gessert *et al.*, 2020; Combalia *et al.*, 2019).

Architectural deep structure and advanced capabilities to learn features allow it to be occupying a special place among the architectures used in complex diagnostic work where it is necessary to distinguish between the minor details of the view and background noise or normal anatomical variations (McKinney *et al.*, 2020; De Fauw *et al.*, 2018). Hierarchical feature extraction offers ResNet50 the ability to understand and extract both low-level textural information and high-level semantics which are vital in the medical diagnosis process.

Recently, it has been shown that residual connections in ResNet50 allow learning clinically useful features and retain stability during training with as little as 2500 images of medical images (Raghu *et al.*, 2019; Azizi *et al.*, 2021). The skip connections can be also used to maintain desirable features in the forward processing and to flow gradients smoothly in the backward CNN operation, and hence learnable to represent medical images patterns well.

Clinical validation efforts indicate it is now possible to attain performance in vascular and other domains on par with or even surpassing human expert level performance using ResNet50-based systems on diabetic retinopathy screening, as well as chest X-ray and skin cancer diagnosis (Gulshan *et al.*, 2016; Rajpurkar *et al.*, 2017; Esteva *et al.*, 2017).

Nonetheless, such achievements have so far been recorded mainly on a controlled background involving well-selected datasets.

Computational Properties and Requirements on Resources

ResNet50 is a computationally very demanding network, and it has around 25.6 million parameters, which makes it quite expensive in terms of computations when it comes to both training and inference processes (He *et al.*, 2016; Bianco *et al.*, 2018). Complexity and depth of the network require extensive memory to train overall, and most often high-end GPU hardware with a minimum of 8GB of memory would be needed to process a batch.

The inference time of ResNet50 also differs greatly depending on hardware setup and optimization methods, and on newer GPU devices with 10-50 milliseconds per image between the best and worst results, and hundreds of milliseconds on less optimized CPU-based inference (Bianco *et al.*, 2018; Canziani, Paszke and Culurciello, 2017). Such timing properties are very relevant to clinical application especially in real time or throughput applications.

Computational requirements of the architecture also include the energy requirements which may be high in resource-limited setting (Strubell *et al.*, 2019; Thompson *et al.*, 2022). Due to the huge computational resources needed per forward pass, the deep architecture is more costly to run than the architecture designed to be deployed on mobiles or edges.

The process of training ResNet50 usually takes a number of hours and a number of days in terms of the size of the dataset and the configuration of the hardware, consuming a significant amount of energy in terms of large-scale projects of medical AI development (Patterson *et al.*, 2021; Schwartz *et al.*, 2020). The computational demands bear relevance

to the accessibility of AI development in a resource-constrained environment and the impact of medical AI research and science on the environment.

2.2.2 Comparative Architectural Analysis

Complexity and Efficiency of the Parameters and Model

The three architectures indicate their varying views of tradeoffs between model simplicity and performance, which has notably important consequences to the potential deployability of the model and clinical usefulness (Canziani, Paszke and Culurciello, 2017; Bianco *et al.*, 2018). ResNet50 has the large representational capacity with 25.6 million parameters and demands massive computational power; thus, no doubt, it would only suit high-performance contexts where the restriction of computer power is minimal.

The 3.5 million parameter MobileNetV2 exploits architectural inventions to make major leaps in efficiency with applications in the real world in resource-constrained environments without suffering a significant loss in performance (Sandler *et al.*, 2018). MobileNetV2 would be especially useful at mobile health operations and point-of-care deployment, where computing power is very scarce.

The middle ground, EfficientNetB0, has 5.3 million parameters and is the model that is carefully optimised, achieving better accuracy than the MobileNetV2, but it is also much more efficient than the ResNet50 (Tan and Le, 2019). Such a trade-off renders EfficientNetB0 favourable to clinical practice when performance and efficiency are a factor.

The differences in the parameter efficiencies reflect directly on the computational resources, memory usage, and energy consumption, and these factors determine the cost and ease of deploying the systems in various healthcare environments in a significant way (Thompson *et al.*, 2022; Strubell *et al.*, 2019).

Medical relevance and Feature Learning Capabilities

The possibilities of learning the medically relevant features and patterns vary under different architectural approaches (Raghu *et al.*, 2019; Azizi *et al.*, 2021). The deep residual network architecture ResNet50 can learn complicated hierarchical representations that could be vital in differentiating subtle pathological variations against normal anatomical variations. The skip connections aid in maintaining fine details and allow the high-level semantic interpreting.

MobileNetV2 depthwise separable convolution is capable of capturing spatial features with high accuracies, and possibly are unable to learn complex cross-channel correlations that are relevant to medical diagnoses (Sandler *et al.*, 2018; Chollet, 2017). Nonetheless, the inverted residual architecture will solve some of these drawbacks since optimal representation capacity becomes amplified in areas that matter the most.

The model has a compound scaling (optimal scaling to reach efficiencies) and squeeze-and-excitation modules, which makes it an attractive solution to medical applications due to efficient learning of complex features with low computational cost (Tan and Le, 2019; Hu, Shen and Sun, 2018). The optimisation of architecture is done in a systematic manner, which allows successful use of representational capacity at every level of abstraction, i.e. low-level texture analysis and up to the high-level semantic comprehension.

Squeeze-and-excitation modules used in attention mechanisms in EfficientNetB0 offer an explicit channel-wise process that can be extremely useful in medical scenarios when some of the feature channels are of stronger diagnostic interest (Wang *et al.*, 2020) (Woo *et al.*, 2018). This focus ability can assist the network to concentrate on clinically significant patterns as well as eliminate the irrelevant background information.

Clinical Deployment questions

The choice of architecture is an intricate interplay of performance, efficiency, and feasibility of deployment that has to be examined carefully in clinical setting (McKinney *et al.*, 2020; De Fauw *et al.*, 2018). The high representational capacity of ResNet50 could be defended in applications in which the diagnostic situation is highly stakes and numerically demanding computational resource is not a problem.

The MobileNetV2 is the best option in instances that the limitations to computations are extreme, or a wide deployment application archetype (Udrea *et al.*, 2020; Finnane *et al.*, 2017). The fact that the architecture is capable of running on commodity mobile hardware may potentially be used to realize diagnostic applications in resource-constrained environments that cannot provide an environment to run-on high-performance computing systems.

The balanced nature of EfficientNetB0 suits several clinical applications because of the desirable accuracy and efficiency associated with it (Tan and Le, 2019). Scalability of architecture offered by compound scaling offers further flexibility to adjust to changing deployment scenarios and performance demands over time.

The complexity of the clinical integration, such as the requirements of real-time processing, batch processing and integration with existing healthcare IT infrastructure, has a considerable impact on choices of architecture (Chen, Liu and Peng, 2019; Sendak *et al.*, 2020). Each architecture has various strengths and constraints to the integration of clinical workflows.

2.3 Current Research Limitations

2.3.1 Methodological Fragmentation and Comparability Issues

The existing situation in CNN architecture research devoted to the area of melanoma detection is characterised by a high degree of methodological fragmentation that

adversely affects the practical translation and evidence-based decision-making processes in clinical implementation (Varoquaux and Cheplygina, 2022; Winkler *et al.*, 2019). The fragmentation has been observed in several aspects of research design such as the selection of dataset and its preprocessing, the evaluation steps, as well as the strategies of statistical analysis.

Selection Bias and by Dataset Inconsistencies

Standardised datasets and assessment procedures of various studies are one of the greatest problems of the present-day research area (Tschandl *et al.*, 2020; Combalia *et al.*, 2019). Scholars most commonly use various datasets, including privately owned clinical databases to openly available databases, and the comparison of their findings is almost impossible. Although the ISIC datasets, HAM10000 and SIIM-ISIC challenges are significant benchmarks in the direction of standardisation, a large proportion of the studies still use institution-specific datasets or a mix of two or more sources, exposing confounding variables that can one way or another erase any important conclusion.

The process of selecting a dataset is mostly influenced by institutional bias and availability and not the clinical relevance or generalisability (Larrazabal *et al.*, 2020; Oakden-Rayner *et al.*, 2020). Investigations that are done in academic hospitals can use data that are different in terms of demographic composition, imaging protocols, and diagnostic criteria compared to the datasets used in community hospital studies, so there may be less validity in estimating performance across clinical practice environments.

Moreover, the data compilations of the studies are highly uneven, as different methods of class imbalance, demographic representation, and image quality standardisation are reflected (Buda, Maki and Mazurowski, 2018; Johnson *et al.*, 2019). In some studies random sampling was provided without clinical factor given consideration, and in another study, there are some studies that try to balance the demographical and

clinical variables by failing to record the selection criteria used and the impacts pertaining to result generalizability.

Because there are no standardised datasets curation and quality assessment procedures, the architecture comparison may be conducted over essentially different dataset properties, thus the reliability of architecture comparison is questionable (Kather *et al.*, 2019; Zech *et al.*, 2018). All this issue is further complicated by the fact that the nature of datasets and preprocessing steps and quality control mechanisms are underreported.

Variations of Preprocessing Protocols

The procedures related to the preprocessing of medical images before the training of the CNN can drastically influence their performance, but no systematic studies have been done on their effectiveness (Perez *et al.*, 2018; Bissoto *et al.*, 2018). Various studies are based on different procedures regarding image normalisation, resizing, augmentation, and quality filtering, which makes finding a common ground hardly accessible.

Colour normalisation is a highly problematic field, because dermoscopic images have color features that can vary critically according to imaging devices, skin light, and capture parameters (Barata *et al.*, 2014; Celebi *et al.*, 2015). There are studies that incorporate elaborate methods of colour normalisation, which have been specifically developed in dermoscopic images, and there are the so-called generic methods due to being intended towards natural images. It is infrequent to have systematic assessment and documentation of the impacts of these decisions on the performance of the model.

Data augmentation can be very basic (e.g., basic rotational or scale transformations), or complex (e.g., domain-specific based aimed at simulating realistic variations of dermoscopic images) (Perez *et al.*, 2018; Shorten and Khoshgoftaar, 2019). The selection of the augmentation technique suffers a similar problem: it is based on

intuition or precedent, instead of a systematic assessment of the effect on the robustness and generalizability of the model.

Lack of documentation in most studies allows many of them to contain critical preprocessing decisions which are not reported or are not documented adequately to reproduce them (Gundersen and Kjensmo, 2018; Hutson, 2018). Such a documentation deficiency results in the fact that it is not possible to tell whether the better performance is attributed to architectural advantage or preprocessing optimisation.

Evaluation Metric inconsistencies

This is possibly the most crucial weakness of those studies available as there is not a consistent use of the evaluation measures that fail to target measures of clinically relevant performance characteristics (Davis and Goadrich, 2006; Saito and Rehmsmeier, 2015). Most metrics take an exclusive look at overall accuracy; however, this does not specifically reflect accuracy in an imbalanced medical dataset, where the performance of the minority classes can be of more clinical value than overall classification accuracy.

Various trials utilize varied base overview measurements, some of the accuracy of measurements, others, fidelity or specificity, and others, area under the curve (AUC) measurements (Tschandl *et al.*, 2020; Combalia *et al.*, 2019). This difference renders direct comparison of studies very hard and renders any meaningful meta-analyses in order to bring evidence-based clinical decision making.

Furthermore, numerous research studies do not include statistics of the comparisons of the performance results and multiple studies statistically compare the results without including confidence intervals or significance testing (Demšar, 2006; Benavoli *et al.*, 2017). Otherwise, without statistical analysis, it is not possible to conclude whether the performance differences are statistically significant or simply the result of chance, and therefore, one could not make responsible conclusions of the differences between them.

It is a so-called multiple comparison problem whereby when testing multiple architectures and multiple metrics the risk of Type I error becomes big (Benjamini and Hochberg, 1995; Holm, 1979). This statistical ignorance is capable of creating fallacies concerning the architectural greatness due to the flukes, instead of true performance benefits.

2.3.2 Limited Integration of Interpretability Assessment

High accuracy is not the only necessity in the successful implementation of AI-assisted diagnostic systems in clinical practice but the possibility of healthcare providers to explain and justify model-made decisions (Holzinger *et al.*, 2017; Rudin, 2019). Nevertheless, existing studies in CNN design in melanoma detection prove poor incorporation of interpretability metrics evaluation, which is a big loophole encountered between technical design and clinical implementation specifications.

Weak Clinical interpretability emphasis

Most present work aims to pay little attention to indicators of performance, but to the issue of how clinical practitioners may interpret and prove the decisions made by the model (Adadi and Berrada, 2018; Tjoa and Guan, 2021). Although interpretability tools like Grad-CAM, saliency maps, or attention visualisation have been developed by researchers to investigate the CNN decision-making process, they are not always employed in the studies of CNN-based melanoma detection systems and lack the rigour of the analysis.

In the cases of interpretability analyses, clinical professionals are not often included in the medical relevance of identified features or attention patterns (Chen and Asch, 2017). Measures of technical interpretability can also determine regions of the image that play key roles in an image, in modelling decisions, but these regions may or may not be valid

diagnostically-informative features, they could be spurious correlates that interfere with clinical reliability.

Moreover, existing interpretability studies routinely use visualisation techniques that are designed to work with general computer vision problems without linking them to the specific needs and limitations of medical diagnosis (Samek and Müller, 2019; Montavon, Samek and Müller, 2018). Clinical interpretability cannot be simply achieved by localising image areas of significance, but it has to be through explanations that can be compatible with accepted diagnostic standards and clinical expertise.

A mismatch between the technical interpretability features and clinical needs is specifically problematic as such explanations should be available and must be able to be verified by healthcare providers on the basis of their clinical experience and used to make an informed decision concerning the AI recommendation (Cabitza *et al.*, 2017; Chen *et al.*, 2020).

Absence of a Repetitious Assessment of Standardised Interpretability

No established process or metrics are developed to measure the traits of interpretability, so it is not possible to compare various architectures in a more systematic way regarding explainability (Doshi-Velez and Kim, 2017; Lipton, 2018). Various studies use different interpretability techniques, but do not have a systematic evaluation of these approaches in terms of effectiveness and reliability.

Although several methods of visualisation are developed, there is no clear agreement on how to quantify the quality of explanations such techniques give or even their clinical significance (Samek, Wiegand and Müller, 2017; Ancona *et al.*, 2017). Other single methods of interpretability used have no validation; on the contrary, others use many methods and have not evaluated their adherence or come up with combined methods that would give a better explanation.

There is no way of performing systematic comparisons of explainability properties of CNN architectures (just as there is no systematic way of knowing whether there are any systematic comparisons ever performed), or of preferentially choosing an architecture based on its properties with respect to interpretability (Murdoch *et al.*, 2019; Rudin, 2019).

Medical applications are especially challenging in terms of such a standardisation gap because their interpretability requirements can be higher than in other fields due to safety, regulatory, and professional responsibility requirements (Holzinger *et al.*, 2017; Amann *et al.*, 2020).

Clinical Gap of Validation

Most critically, the contemporary interpretability research tends to fail to give proper clinical validation of explanations given by various approaches (DeGrave, Janizek and Lee, 2021). Statistically significant patterns in model attention or feature attribution may be detected by technical measures of interpretability, but their diagnostic value is unknown unless they are reviewed by a clinical expert.

There are few studies with clinical validation, which in most cases comprise a small population of clinical specialists or do not have a strict protocol regarding inter-rater reliability and inter-interpretability reliability assessment (Chen *et al.*, 2020; Cabitza *et al.*, 2017). This is a serious disadvantage to the applicative utility of interpretability research and retards the growth of explainable AI systems capable of clinical use.

Systematic visualisation of interpretability should be evaluated by board-certified dermatologists to know whether it is clinically relevant, and whether there are some problems with the way models make decisions, which is never done with sufficient rigor (Winkler *et al.*, 2019; Haenssle *et al.*, 2018).

As with any clinical validation gap, the implications of such modern research into interpretability lack in real-world applicability to the clinical situation (Topol, 2019; Yu *et al.*, 2018).

2.3.3 Statistical Rigour and Reproducibility Challenges

Recent and noteworthy studies on CNNs architectures in the context of melanoma detection and diagnostics are often deficient in statistical rigour and reproducibility required to make scientific advancements and practice-relevant clinical decisions (Varoquaux and Cheplygina, 2022; Gundersen and Kjensmo, 2018). Such obstacles make research results questionable and hinder technological breakthroughs to be used in line with practice.

Inappropriate Statistical Analysis

Most of the research on melanoma detection use inappropriate statistical analysis, which decrease the credibility and validity of the results (Demšar, 2006; Benavoli *et al.*, 2017). Point estimates of performance metrics are often published without suitable confidence intervals, and so the statistical significance of architectural differences or reliability of performance claims cannot be estimated.

Experimentation using statistical tests may be difficult when testing model performance, and in many studies, statistical tests used do not consider the paired nature of architectural comparisons or the problem of multiple comparisons, providing statistical tests comparing many architectures at a time (Benjamini and Hochberg, 1995; Holm, 1979). A single-statistic correction of such multiple comparisons may result in an excess of the Type I error rate and false claims of significance.

The size of samples is hardly ever mentioned or even rationalised, although the sample size calculation is required to achieve the right amount of power to differentiate architecture versions of clinically significant differences (Cohen, 1998; Faul *et al.*, 2007).

Most studies depend on a convenience sample consisting of readily available data as opposed to principled sample size selection that is derived by the effect size and the preferred statistical power.

The complex nature of medical AI evaluation makes the problem of statistical analysis worse, and traditional statistical techniques might not be suitable to handling hierarchical data structures, clustering effects, and non-independence-related issues existing in the medical datasets (Krstajic *et al.*, 2014; Roberts *et al.*, 2021).

Lack of documentation and reproducibility

Reproducibility in scientific research has become an acute problem with deep learning research especially since complicated software settings, stochastic training algorithms, and commercial data modes may render reproduction of outcomes a near-impossibility (Gundersen and Kjensmo, 2018; Hutson, 2018). A great number of studies do not offer enough detail of experiments, hyperparameter selection, or implementation decisions to permit reliable replication.

The level to which the code is available and how good it has been documented differs considerably between studies with many having inadequate documentation to reproduce their work (Stodden, Seiler and Ma, 2018; Peng, 2011). Even with code present, the code is usually not well documented or uses dependencies that are undocumented to the point of being non-reproducible.

Reproducibility of deep learning experiments depends on random seed management and control of stochastic processes, although these aspects are discussed and well-regulated too seldomly (Henderson *et al.*, 2018; Bouthillier *et al.*, 2021). Random initialisation, data shuffling, and stochastic optimisation procedures used in training can cause significant variation across running cycles, but are frequently not well-controlled experimentally and reported.

Constraints of External Validation

A great deal of existing studies do not provide enough external validation between different datasets, populations, and clinical contexts to show that research can be generalizable (Keane *et al.*, 2018; Liu *et al.*, 2020). There is also limited cross-validation in the studies that were conducted, relying only on one dataset or sometimes minimal validation techniques, which do not necessarily evaluate the cross-validation in varied conditions, as in clinical practice.

The methods of validation applied are not always representative of the real-life clinical implementation setting and rely on an artificial test set, which does not always reflect the shifts of distributions, the change of quality, and the demographic differentiation in the real-life clinical setting (Zech *et al.*, 2018; Larrazabal *et al.*, 2020). This drawback hinders the possibility of forecasting how the results of the research are likely to be applied to clinical practice.

The validation of geography and demography is highly restricted, and most of the studies base their work on the data of certain institutions or areas, which do not reflect the human population all over the world (Larrazabal *et al.*, 2020; Seyyed-Kalantari *et al.*, 2021). Recent studies have not clearly defined the performance of CNN architectures in various demographic populations, skin types, and imaging situations.

Temporal validation, the testing of models over time series data collected at different time points, is also seldom performed even though it is relevant to know how models might perform over time (Davis *et al.*, 2017; Subbaswamy and Saria, 2020). This shortcoming is especially problematic when focusing on clinical uses in which model performance can deteriorate with time since imaging equipment or protocols are altered or the patient population changes.

All of these methodological shortcomings aggregate to compromise the reliability and clinical relevance of ongoing studies, which makes them a source of concern because substantial experimental design improvements, statistical analysis, and validation models are essential to the research concerning the CNNs architectures in detecting melanoma.

CHAPTER III: METHODOLOGY

3.1 Research Design Framework

3.1.1 Philosophical Foundation and Research Paradigm

The research paradigm embraced by the study is a positivist paradigm, and it utilises a quantitative experimental paradigm to come up with objective and reproducible evidence of CNN architecture performance in detecting melanoma (Creswell and Creswell, 2017; Johnson and Onwuegbuzie, 2004). This school of thought admits that to achieve substantive growth in medical AI, systematic, empirical testing is necessary, which can consistently assist in clinical decision-making. The study adheres to a deductive reasoning, basing itself on the existing CNN theory to generate testable hypotheses on the effectiveness of architectures in medical imaging (Bryman, 2016; Saunders, Lewis and Thornhill, 2009)), which fits the practical objective of generating evidence-based implementation findings as opposed to building new theory.

The design of the research focuses on the methodological rigor and experimental control so that the observed differences of performance are due to architectural differences instead of confounding factors (Shadish, Cook and Campbell, 2002; Campbell and Stanley, 1963). Through controlled comparisons, the study actively goes to the root of significant methodological disparities in the existing literature, in which comparatively meaningless experimental variation has precluded meaningful benchmarking of architecture. This method makes the results more authoritative, and the input to the field of medical AI assessment more convincing.

Epistemological Considerations

The epistemological basis of this research is that the credible information on the performance of AI systems should be based on systematic and empirical assessment instead

of theoretical speculations or anecdotal observations (Popper, 2002; Kuhn, 1962). Since architectural design decisions, training protocols and dataset specifications are highly coupled, a careful experimentation is necessary in generating plausible results. Although the research admits that quantitative measures cannot comprehensively address the issues of clinical interpretability and workflow integration as significant areas of qualitative consideration in medical AI assessment (Tashakkori and Teddlie, 2010; Creswell and Clark, 2017), it nevertheless holds that the quantitative comparison is a fundamental aspect of evidence-based decision-making, which is further supplemented by qualitative evidence. The fact that it includes the confidence intervals and uncertainty quantification is also an additional strength of the research design, as it considers the natural variability of machine learning experiments and does not allow concluding only on unreliable point estimates (Wasserstein and Lazar, 2016; McShane *et al.*, 2019).

3.1.2 Experimental Design Framework

Controlled Comparison Methodology

The research follows a controlled comparative experimental design that has reduced the confounding variables to the extent that all other variables have been standardised and only the architectural features of interest have been altered (Montgomery, 2017; Box, Hunter and Hunter, 2005). Architectures are completely standardised in dataset partitioning, preprocessing, training, and evaluation metrics in order to make fair comparisons (Hastie *et al.*, 2009; Bishop, 2006). Randomisation and multiple random seed trials minimise the reliance of special data structures and improve the results stability (Fisher, 1935). Also, the experimental environment imposes high level of hardware, software environment and parameter records of a study to guarantee reproducibility and verification of results (Peng, 2011; Stodden, Seiler and Ma, 2018).

Multi-Dimensional Evaluation Framework

The assessment system includes several areas that are necessary to implement the use of medical devices, such as technical performance, practicality, interpretability, and statistical reliability (Liu *et al.*, 2019; Nagendran *et al.*, 2020). The clinically relevant metrics used to measure performance include sensitivity and specificity where the performance of minority classes is given special attention since melanoma detection is a critical aspect (Davis and Goadrich, 2006; Saito and Rehmsmeier, 2015). Computational evaluation determines the use of memory, energy, and general resource requirements in hardware platforms (Canziani, Paszke and Culurciello, 2017; Bianco *et al.*, 2018). The interpretability analysis uses systematic visualisation methods to examine the matching of the model decision-making with medical reasoning (Selvaraju *et al.*, 2017; Adadi and Berrada, 2018).

Statistical Analysis Framework

The assessment system includes several areas that are necessary to implement the use of medical devices, such as technical performance, practicality, interpretability, and statistical reliability (Liu *et al.*, 2019; Nagendran *et al.*, 2020). The clinically relevant metrics used to measure performance include sensitivity and specificity where the performance of minority classes is given special attention, since melanoma detection is a critical aspect (Davis and Goadrich, 2006; Saito and Rehmsmeier, 2015). Computational evaluation determines the use of memory, energy, and general resource requirements in hardware platforms (Canziani, Paszke and Culurciello, 2017; Bianco *et al.*, 2018). The interpretability analysis uses systematic visualisation methods to examine the matching of the model decision-making with medical reasoning (Selvaraju *et al.*, 2017; Adadi and Berrada, 2018).

3.1.3 Quality Assurance and Validation Procedures

Experimental Validity Measures

The validity measures in the study are rigorous so as to provide reliability, generalisability and clinical relevance of findings (Campbell and Stanley, 1963; Cook and Campbell, 1979). The controls of experiments and randomisation are applied to enhance internal validity by reducing effects of confounding factors and external validity is enhanced by the subgroups to determine performance consistency between demographics and lesion types across the use of one main dataset (Shadish, Cook and Campbell, 2002; Rothwell, 2005). Construct validity is upheld because the evaluation metrics are chosen based on the clinical performance requirements as opposed to technical indicators only (Cronbach and Meehl, 1955; Messick, 1995). The appropriate sample sizes, high levels of control of Type I/II errors, multiple comparisons correction and presentation of confidence intervals all support statistical conclusion validity (Cook and Campbell, 1979; Maxwell, Delaney and Kelley, 2017).

Reproducibility and Documentation Standards

The research complies with high standards of reproducibility, that is backed by detailed records of all the experimental procedures, implementation and analysis choices (Stodden, Seiler and Ma, 2018; Nosek *et al.*, 2015). Hyperparameters, training behaviour, and the entire hardware-software environment can be logged, which guarantees that experiments are exactly repeatable (Henderson *et al.*, 2018; Bouthillier *et al.*, 2021). Methodological decisions are also justified by decision logs, which document alternatives eliminated as a result of sensitivity analysis and encourage transparency in the interpretation (Kapoor and Narayanan, 2023; D'Amour *et al.*, 2020). Code and configuration changes are tracked by version control systems, which makes it possible to

trace provenance fully and allows a researcher to understand how modifications affect the results (Ram, 2013; Wilson *et al.*, 2017).

Bias Mitigation Procedures

The study adheres to systematic approaches in order to reduce bias and mitigate numerous threats to research integrity (Hernán and Robins, 2020; Greenland *et al.*, 2016). Structured dataset partitioning of data to control selection bias is to ensure representative samples, and multiple independent evaluation methods are used to control confirmation bias (Nosek *et al.*, 2018; Simmons, Nelson and Simonsohn, 2011). Transparency in recording all the findings, such as non-significant or adverse ones, minimises reporting bias (Dwan *et al.*, 2008; Cooper, Hedges and Valentine, 2009). Clinical relevance of evaluation metrics, strict statistical methods, and systematic controls of confounding variables also prevent the occurrence of analytical bias, and performance measurements are associated with actual clinical demands (Luo *et al.*, 2016; Flach, 2019).

3.1.4 Ethical Considerations and Compliance

Human Subjects Protection

The research involves fully de-identified publicly accessible data and hence does not qualify as Human Subjects Research as stipulated by the federal regulations (Office for Human Research Protections (OHRP), 2025; Intramural Research, 2025). Nevertheless, the study adheres to the standards of responsible medical AI creation, such as secure handling of data, restricted access, and adherence to health information protection requirements (U.S. Department of Health and Human Services, Office for Civil Rights, 2025; Intersoft consulting, 2025). Wider moral aspects like fairness, access, and bias are built in with the help of demographic subgroup performance analyses to determine possible differences (Char *et al.*, 2018; Rajkomar *et al.*, 2018). Additional principles of beneficence and non-maleficence guide the research and help to care about patients and avoid risks

connected with premature or unsuitable clinical application of AI systems (Beauchamp and Childress, 2019; Emanuel, 2000).

Scientific Integrity and Transparency

The research is conducted according to the existing guidelines that facilitate scientific integrity and detailed reporting and ethical research conduct, and the limitations and conflicts of interest is disclosed (Academies *et al.*, 2017; Committee on Publication Ethics (COPE), 2025). Results are shared following AI transparency guidelines such as the complete reporting of negative outcomes, uncertainties, and methodology limitations to prevent the selective publication of research or overstated claims on generalisation (Mongan *et al.*, 2020; Liu *et al.*, 2019). The information on financial disclosures, institutional affiliations, and collaborative agreements is recorded so that the assessment of potential sources of bias could be performed properly (International Committee of Medical Journal Editors (ICMJE), 2025; Drazen *et al.*, 2010). Best-practice guidelines of data-sharing and collaboration do not exclude transparency and give due respect to intellectual property and research integrity (National Science Foundation (NSF), 2025; Wilkinson *et al.*, 2016).

Responsible AI Development

The study design corresponds to the principles of responsible AI development that focus on patient safety, clinical relevance, and fair access and does not involve technical optimisation (Floridi *et al.*, 2018; Jobin, Ienca and Vayena, 2019). The assessment system focuses on clinically significant performance indicators and looks at potential failure modes and their consequences to patient outcomes. The broader societal effects, including alterations to the healthcare access, medical practice, and patient experience in various systems, are also taken into account, particularly the potential of melanoma-detection AI tools to shape the practice and patient experience (Topol, 2019; Yu *et al.*, 2018). The

research inspires further research to deal with the limitations reported and continue on the road of safe and equitable AI implementation, by recognizing the fact that medical AI creates ethical obligations not necessarily aligned with established scientific inquiry (Beam and Kohane, 2018; Char *et al.*, 2018); (Emanuel *et al.*, 2020; Matheny, Whicher and Israni, 2020).

3.2 Secondary Data Collection Procedure

3.2.1 SIIM-ISIC 2020 Dataset Characteristics

The SIIM-ISIC 2020 Melanoma Classification Challenge dataset contains the largest and thoughtfully curated selection of dermoscopic images to date to be used in the research of melanoma detection (Rotemberg *et al.*, 2021; Cassidy *et al.*, 2022). The reason this dataset was identified as the most relevant one within the scope of the current research is because of its superior excellence, the excellent annotation and its substantial and widespread use as the benchmark standard in the field of AI in dermatology.

Data Composition and magnitude

It is based on 33,126 high-resolution dermoscopic images that were acquired under standardised imaging protocols in a variety of institutions, which is an adequate size to train deep neural networks and makes it possible to conduct reliable statistical comparisons (Rotemberg *et al.*, 2021). The variety of clinical settings, patient demographics, and geographic areas makes its clinical heterogeneity more favourable to the generalisability of the results, since the findings are not limited to a single centre (Tschandl *et al.*, 2020). Every photograph is of high quality in terms of resolution, focusing, lighting, and visibility, and artefacted or low-quality photos are filtered out during the curation process to make sure that the photograph is scientifically reliable (Codella *et al.*, 2019). The combination of the size of the dataset, its variety, and its quality results in it being very well-suited in

terms of comparative architectural assessment since it can measure the difference between models in fine detail without being obscured by poor data quality (Tschandl *et al.*, 2020).

Clinical Annotation and Verification

Images are all annotated by board-certified dermatologists, and often have additional dermatopathology expertise, and have histopathological confirmation where it is available, which forms an effective gold standard of supervised learning (Rotemberg *et al.*, 2021). The variability between observers is reduced through standardised diagnostic criteria and cross-institutional quality-control processes, and challenging cases are subject to multiple reviews of experts in order to reach a consensus (Tschandl *et al.*, 2020; Kittler *et al.*, 2002). Other quality assurance measures are senior dermatologists to review the annotations and cross-labelling with available pathology reports and to exclude ambiguous cases that may affect the quality of training (Marchetti *et al.*, 2018; Tschandl *et al.*, 2019). Quality, consistent labels are needed to eliminate potential misleading architectural comparisons because the noise of labels may have a disproportionate influence on model performance and mislead conclusions (Northcutt, Jiang and Chuang, 2021; Zhang *et al.*, 2021).

Characteristics of class distribution and imbalance

The dataset represents the clinical screening setting, with 584 melanoma cases (1.76) and 32,542 benign lesions (98.24), which leads to an imbalance of 56:1, and typically a clinical population (Rotemberg *et al.*, 2021). Such a striking imbalance poses fundamental methodological concerns, because even traditional methods of machine learning can attain artificially high accuracy and yet not satisfy the clinical safety levels (Buda, Maki and Mazurowski, 2018; Johnson and Khoshgoftaar, 2019). The dataset hence offers a useful test bed to check the resilience of CNN architectures to approximate class distributions as well as to determine architectural choices that may be reliable even given

few positive cases (Krawczyk, 2016; He and Garcia, 2009). This imbalance is essential to be analysed because when the ratio is that high conventional algorithms tend to fail to detect the minority-class (Japkowicz and Stephen, 2002; Chawla *et al.*, 2002).

Metadata and Clinical Context

All the images have rich metadata (age of the patient, sex of the patient, location of the lesion, and the clinical situation) that allow analysing the images beyond a simple classification (Rotemberg *et al.*, 2021; Cassidy *et al.*, 2022). The anatomical location metadata is especially helpful to gain insight into diagnostic issues specific to certain body part locations because each site of the body varies in terms of skin texture, hair presence, lesion morphology and imaging artefacts (Fidalgo Barata, Celebi and Marques, 2014; Mendonca *et al.*, 2013). Demographic data can be used to assess the algorithmic fairness, which enables one to determine the differences in performance between groups of patients (Larrazabal *et al.*, 2020; Seyyed-Kalantari *et al.*, 2021). Further development of clinical context lesion characteristics and diagnostic considerations also expand further evaluation by providing more clinically based insights into model behaviour (Marchetti *et al.*, 2018; Tschandl *et al.*, 2019).

3.2.2 Data Quality Assessment and Preprocessing

Assessment of an Image Quality

Quality evaluation of all images was conducted both by automated and hand standards to match the criteria that could have been used to guarantee quality analysis (Fidalgo Barata, Celebi and Marques, 2014; Celebi *et al.*, 2015). Sharpness, illumination uniformity, and all common dermoscopic artefacts in the image (hair or ruler marks) were assessed automatically with the evaluated image being flagged as requiring review or deletion when the score fell below the minimum threshold (Emre Celebi *et al.*, 2013; Garnavi *et al.*, 2011). The presence of sufficient visibility of lesions, the presence of an

appropriate magnification, and the absence of major obstruction were confirmed by manual clinical review in order to avoid the presence of technically acceptable but clinically unusable pictures (Kittler *et al.*, 2002; Argenziano *et al.*, 2003). Recording quality problems gave an idea of realistic imaging problems and assisted in determining proper clinical-level quality controls (Fidalgo Barata, Celebi and Marques, 2014; Celebi *et al.*, 2015).

Procedures of Standardization and Normalization

The use of standardisation procedures was to remove any non-clinical variability based on imaging hardware and acquisition settings and maintain the diagnostically relevant features (Perez *et al.*, 2018; Bissoto *et al.*, 2018). All the images were resized to 224x224 pixels with high-quality interpolation to preserve most important visual information and match the popular network input standards (Keys, 1981; Meijering, 2002). Colour normalisation has fixed light and device differences, statistical with Compatibility using ImageNet standards, so the pre-trained models can be compatible (Barata *et al.*, 2014; Celebi *et al.*, 2015). The pixel-value normalisation also made sure that there was consistency of the input without any change in the intensity relationship that is important to the clinic, which favoured the transfer-learning pipelines (Deng *et al.*, 2010; Russakovsky *et al.*, 2015).

Partitioning Strategy of Data

The stratified random sampling was used to preserve the natural class proportions of the melanoma and the benign lesions between training (70%), validation (15%), and testing (15) subsets (Kohavi and others, 1995; Stone, 1974; Hastie *et al.*, 2009). Stratification created consistency in representing the minority melanoma group and also taken into account the important metadata like location of anatomy and patient demographics to maintain a balance in the subgroups (Japkowicz and Stephen, 2002;

Chawla *et al.*, 2002). The use of fixed random seeds was to be able to guarantee reproducibility and prevent data leakage, to which it was possible to add patient-level stratification to prevent the possibility of lesions in different partitions belonging to the same person (Henderson *et al.*, 2018; Bouthillier *et al.*, 2021). This plan was effective in giving strong training conditions and good evaluation without the bias that may occur due to unequally sampled patients.

Validation and Quality Control

The integrity of the dataset and reliability of the experiment were also guaranteed by quality control procedures by validating proper pre-processing, label consistency, and proper partitioning (Northcutt, Jiang and Chuang, 2021; Zhang *et al.*, 2021). Diagnostic tests indicated that the determination of class distribution, demographic and quality traits were similar over partitions so that performance variation was due to architectural behaviour as opposed to sampling artefacts (Kohavi and others, 1995; Stone, 1974). Preprocessing validation entailed the systematic verifications and representative hand checks to ensure that all the images were subjected to the same transformations without creating distortions or loss of diagnostic features (Perez *et al.*, 2018; Bissoto *et al.*, 2018). The validation of annotations was made to ensure that labels, metadata, and identifiers were consistent during data preparation; hence, all other training and evaluation processes were valid (Northcutt, Jiang and Chuang, 2021).

3.2.3 Data Augmentation and Enhancement

Medical-Domain-specific Augmentation Strategy

The augmentation strategy has been created to add clinically realistic variability to the dermoscopic imaging, without altering melanoma-relevant diagnostic characteristics (Perez *et al.*, 2018; Shorten and Khoshgoftaar, 2019). Natural changes in imaging orientation were geometrically modelled by full-range rotations, flips, and constrained

scaling, which did not change lesion morphology (Simard, Steinkraus and Platt, 2003; Krizhevsky, Sutskever and Hinton, 2012). Manipulated colour values recreated variation in light and camera parameters and preserved diagnostically significant chromatic patterns (Perez *et al.*, 2018; Buslaev *et al.*, 2020). The anatomical variability derived by the introduction of limited deformations, elicited by nature, into the skin topology (Simard, Steinkraus and Platt, 2003; Ronneberger, Fischer and Brox, 2015). In general, the approach made the dataset diversification balanced and severe at clinical validity.

Augmentation Clinical Validation of

Clinical validation tests were used to make sure that augmented images were diagnostically viable and without unrealistic distortions (Fidalgo Barata, Celebi and Marques, 2014; Celebi *et al.*, 2015). The review of experts also proved that the most important lesion structures were maintained during transformations, and automated filters selected augmented samples with artefacts or the absence of diagnostic detail (Bissoto *et al.*, 2018; Pacheco *et al.*, 2020). The systematic augmentation bias was avoided by randomisation of augmentation parameters within medically validated ranges (Cubuk *et al.*, 2020; Lim *et al.*, 2019). The parameter distributions and the effects of their impact on the model training dynamics were further evaluated through statistical monitoring to stipulate that augmentation facilitated meaningful learning (Muller and Hutter, 2021; Fort, Hu and Lakshminarayanan, 2020).

Augmentation Impact Analysis

In impact analysis, the effects of the type and intensity of augmentation on the model accuracy, generalisability, and robustness were studied (Cubuk *et al.*, 2020; Lim *et al.*, 2019). Ablation experiments have revealed changes that have the largest impact on performance enhancement, and such changes should be eliminated to allow eliminating of low-benefit or potentially harmful augmentations (Muller and Hutter, 2021; Fort, Hu and

Lakshminarayanan, 2020). Robustness tests were able to gauge the stability of the model in different imaging conditions and the demographic and lesional subgroups (Hendrycks and Dietterich, 2019; Recht *et al.*, 2019). Computational impact analysis also took into account the overheads in training time and in memory to strike a balance between the complexity of augmentation and the practicality (Cubuk *et al.*, 2020; Muller and Hutter, 2021).

3.2.4 Dataset Management and Version Control

Data Consistency and Reproducibility

Robust data-handling procedures were implemented to ensure data integrity, reproducibility, and controlled access (Wilkinson *et al.*, 2016; Jacobsen *et al.*, 2020). Dataset versioning with cryptographic hashing and checksums safeguarded authenticity, enabling detection of any unintended or unauthorised alterations during processing (Ram, 2013; Wilson *et al.*, 2017). Strict access control protocols restricted dataset availability to authorised personnel, generating audit trails that reinforced privacy protection and regulatory compliance (Technology, 2020). Additionally, geographically distributed backups and multi-layer redundancy supported data continuity and prevented losses caused by technical failures (Hey *et al.*, 2009).

Management and Metadata Documentation

Comprehensive documentation of dataset creation, pre-processing steps, and experimental setups enhanced transparency, reproducibility, and interpretability of results (Goodman *et al.*, 2014; Sandve *et al.*, 2013). Metadata systems recorded detailed information for every image—including lineage, processing history, and quality metrics—enabling nuanced analysis and detection of potential biases (Wilkinson *et al.*, 2016; Jacobsen *et al.*, 2020). Processing logs captured parameters, software versions, and operator actions (Henderson *et al.*, 2018; Bouthillier *et al.*, 2021), while version-controlled

analysis scripts maintained clear provenance of methodological changes and their rationale (Stodden, Seiler and Ma, 2018; Peng, 2011).

Compliance and Security

It implemented security controls and controls, such as data encryption, secure transfer protocols, and monitored access to secure sensitive medical imaging data and to support ethical research use of de-identified datasets. The adherence processes maintained the conformity to the institutional, legal, and professional guidelines that were strengthened by regular audits (Office for Human Research Protections (OHRP), 2025). The combination of research requirements and privacy and storage requirements was achieved using data-retention policies, which provided timelines on how data can be preserved and safely disposed. Partnerships between nations also ensured that collaboration across borders was performed with respect to the jurisdictional privacy standards (Dove, Knoppers and Zawati, 2014).

3.3 Data Preprocessing and Architecture Implementation

3.3.1 Comprehensive Preprocessing Pipeline

Standardising Procedures of Images

Pre-processing pipeline introduces systematic standardisation of eliminating technical differences that might bias architectural comparisons without affecting clinically relevant information to detect melanoma (Perez *et al.*, 2018; Barata *et al.*, 2014). Resolution checks, gradient-based sharpness measures, and histogram-based illumination analysis are automated screening measures of low-quality or clinically inappropriate images (Emre Celebi *et al.*, 2013; Garnavi *et al.*, 2011). The images are all resized to 224x224 by using bicubic interpolation in order to preserve diagnostic information and computational efficiency, which is in line with ImageNet-based workflows (Keys, 1981; Meijering, 2002). Centre-cropping is done to preserve aspect-ratio and resize to prevent

distortion of lesion morphology, which is essential to diagnostic integrity (Wang *et al.*, 2018).

Enhanced Colour Normalisation

Colour normalisation deals with actual hardware and lighting differences that can confound the comparison of models (Barata *et al.*, 2014; Celebi *et al.*, 2015). ImageNet means and standard deviation are employed in standardisation to ensure compatibility with the pre-trained networks, and to examine the relationships between colours in a diagnostic manner (Deng *et al.*, 2010; Russakovsky *et al.*, 2015). Selective adaptive histogram equalisation is only used to excessively bright or dark images so as to enhance contrast without interfering with the natural pigmentation patterns (Pizer *et al.*, 1987; Zuiderveld, 1994). Colour-space Dataset level analysis determines systematic colour bias or outliers, indicating unusually coloured images to be reviewed further (Fidalgo Barata, Celebi and Marques, 2014; Celebi *et al.*, 2015).

Processing Intensity of Pixel

Pixel-intensity normalisation transforms the 0255 values of a raw image into a floating-point 01 range that is used by contemporary neural networks, without changing relative intensity differences that are useful in diagnosis (Goodfellow *et al.*, 2016; LeCun, Bengio and Hinton, 2015). Adaptive contrast enhancement is used to emphasise details of fine lesions and prevent over-enhancement to form artifacts (Pizer *et al.*, 1987; Stark, 2000). Technical noise can be eliminated using noise-reduction filters e.g. edge-preserving smoothing, which preserve the lesion edges and texture patterns (Buades, Coll and Morel, 2005; Dabov *et al.*, 2007). There is also dynamic-range optimisation which uses maximum possible intensity variation across images by normalising the representations and enhancing the consistency of feature extraction (Reinhard *et al.*, 2002; Durand and Dorsey, 2002).

3.3.2 CNN Architecture Implementation

Implementation details EfficientNetB0

EfficientNetB0 has 5.3M parameters and adheres to the principle of compound scaling, it simultaneously scales the network depth, width, and input resolution (Tan and Le, 2019). To provide benchmark compatibility and successful transfer learning, the official implementation of TensorFlow / Keras with ImageNet pre-trained weights is employed. The network consists of seven stages, which are constructed based on MBConv modules, depth wise-separable convolutions, squeeze-and-excitation units, and skip connections, which allow it to extract features efficiently (Tan and Le, 2019)(Sandler *et al.*, 2018). Original coefficients of scaling of compounds ($a=1.2$, $b=1.1$ and $c=1.15$) are kept because they provide good balance between accuracy and efficiency. Transfer learning initially freezes early layers and then unfreezes the deep layers to learn dermoscopic patterns, and the final classifier is substituted with a binary head to identify melanoma (Yosinski *et al.*, 2014; Raghu *et al.*, 2019).

ResNet50 Implementation Configuration

ResNet50 is structured around the common 50-layer, 25.6M-parameter model and utilises ImageNet pre-trained weights in the official version of Keras (He *et al.*, 2016). It has four residual phases with increasingly lower spatial resolutions and increasing channel expansions (56×56 to 7×7 ; 64–512 channels). Each stage has bottleneck residual blocks with 1×1 , 3×3 and 1×1 convolutions computed with feature extraction that is computationally efficient. The Skip connections also facilitate the model to learn the residual mappings, which reduces the problem of vanishing gradient and permits the model to have a deeper hierarchy of representation (He *et al.*, 2016; Veit, Wilber and Belongie, 2016). Post convolutional batch normalisation stabilises training and allows higher

learning rates in addition to adjusting to medical imaging properties (Ioffe and Szegedy, 2015; Santurkar *et al.*, 2018).

MobileNetV2 Implementation Specifications

MobileNetV2 is deployed with its effective 3.5M parameters architecture with ImageNet pre-trained weights (Sandler *et al.*, 2018; Howard *et al.*, 2017). Depthwise separable convolutional networks break down conventional convolutional networks into depthwise and pointwise convolutional networks, which can reduce computation by many orders of magnitude without affecting feature extraction ability (Chollet, 2017; Kaiser, Gomez and Chollet, 2017). The network accommodates the inverted residual blocks and increases the channel dimension with pre-depthwise convolutions expansion factor of 6 is the best combination of representational power and efficiency (Sandler *et al.*, 2018; Ma *et al.*, 2018). Linear bottlenecks eliminate nonlinear activations in compressed spaces to ensure no information is lost in ReLU operations in low-dimensional features (Sandler *et al.*, 2018; J. Zhang *et al.*, 2018).

3.3.3 Training Protocol Standardisation

Optimization Configuration

All models are trained in a standardised manner with the Adam optimiser with tuned hyperparameters in order to stabilise convergence (Kingma and Ba, 2017; Reddi, Kale and Kumar, 2019). The learning rate of 0.001 using exponential decay (0.96 per every 5 epochs) facilitates the smooth optimisation. A batch of 32 is calculated with sufficient computational constraints and gradient stability (Smith *et al.*, 2018). Regularisation of the weight decays avoids overfitting without incurring representational capacity losses (Loshchilov and Hutter, 2019). Gradient clipping (max-norm 1.0) stabilises training on underclass datasets and prevents exploding gradients (Pascanu, Mikolov and Bengio, 2013; Zhang *et al.*, 2019).

Transfer Learning Strategy

ImageNet models are reconfigured via transfer learning that does not negatively affect the performance of low-level feature extractions (Yosinski *et al.*, 2014; Raghu *et al.*, 2019). The initial step is to freeze all the convolutional layers and to train the classification head on 1015 epochs with a learning rate that increases (0.01) so as to transform the medical labels (Donahue *et al.*, 2014; Girshick *et al.*, 2014). A larger unfreezing step is followed with a subsequent step of adjusting deeper layers with lower learning rates (0.001 - 0.0001) to avoid catastrophic forgetting (Howard and Ruder, 2018; Peters *et al.*, 2018). The difference of the rates of learning in the layers also assists to offer a balance between the memorisation of the general features and the customisation of the domain (Chronopoulou, Baziotis and Potamianos, 2019).

Training, Monitoring and Early Stopping

Continuous training is monitored using loss, accuracy, sensitivity, specificity and AUC-ROC to ensure that learning is stable (Caruana, Lawrence and Giles, 2001; Prechelt, 1998). Early stopping is done through the use of validation AUC-ROC (patience 15, min improvement 0.001) to prevent overfitting and maximise learning (Prechelt, 1998). Overfitting, underfitting or training instability is determined by learning curve analysis (Figueroa *et al.*, 2012). Gradient norms, activation statistics, and layer-wise learning rates are kept in order to obtain a diagnostic insight (Hestness *et al.*, 2017). To be able to recover them and compare their performance, checkpointing captures a number of model states (Goyal *et al.*, 2018; You, Gitman and Ginsburg, 2017).

3.3.4 Experimental Control and Reproducibility

Random Seed Management

Random seed control guarantees the consistency of the experimental runs and allows to determine the variability of the performance reliably (Henderson *et al.*, 2018;

Bouthillier *et al.*, 2021). Random number generators (NumPy, TensorFlow, Python) are all initialised using fixed seeds to ensure similar conditions of comparison across architectures. Five seeds are used to train each model in order to measure robustness, measure performance ranges, and allow testing significance (D’Amour *et al.*, 2020; Kapoor and Narayanan, 2023). The shuffling of the data between the runs is also comparable and can be used to analyse the effects of the order of training sensitivity and maintain experiment fairness (Fisher, 1935). Initialisation of weights is performed following architecture-specific schemes: He, which applies to ReLU layers, and Glorot, which applies to others, as well as with fixed seeds to ensure the same starting states (Glorot and Bengio, 2010; He *et al.*, 2015).

Hardware and Software Standardisation

To eliminate the differences in performance that could occur due to differences in computations, all the experiments are conducted with the same hardware, namely, NVIDIA RTX 3080 (10GB), Intel i7-9700K, and 32GB RAM (Henderson *et al.*, 2018; D’Amour *et al.*, 2020). The versions of software are highly monitored, such as TensorFlow 2.8, CUDA 11.2, Python 3.8, and all the necessary libraries, which are kept in a dedicated virtual environment to ensure reproducibility (Stodden, Seiler and Ma, 2018; Peng, 2011). Regular collection of garbage helps to avoid fragmentation of memory, and there is the stability of the runtime conditions (Henderson *et al.*, 2018). A constant check of CPU/GPU usage, RAM usages, and storage I/O will guarantee hardware stability and identify performance anomalies (Thompson *et al.*, 2023; Patterson *et al.*, 2021).

Experimental Logging and Documentation

There is a detailed logging system that captures hyperparameters, training dynamics, performance metrics, and resource usage, among others to achieve transparency and reproducibility (Stodden, Seiler and Ma, 2018; Nosek *et al.*, 2015). Configuration

management is very specific and tracked on model settings, pre-processing protocols, and evaluation protocols, allowing the precise replication and comparison of experimental conditions (Wilson *et al.*, 2017; Ram, 2013). Structured exception logging allows all errors and unsuccessful executions to be logged, which facilitates systematised debugging and systematic refinement of methods (Henderson *et al.*, 2018; D’Amour *et al.*, 2020). Git based version control tracks amendments in code, configurations and documentation offering complete tracking of experimental evolution (Stodden, Seiler and Ma, 2018; Wilson *et al.*, 2017).

3.4 Primary Data Collection Procedure

3.4.1 Overview of Data Collection Approach

The information was collected in the form of a structured questionnaire in the online survey delivered by Google Forms and sent to doctors engaged in dermatology, general medicine, diagnostics, and other spheres. The survey was done to collect the quantitative answer on the perception, attitudes, and acceptance of automated melanoma detection systems of participants based on the concepts of UTAUT and trust-ethics. Convenience sampling was used to get 100 valid responses and was chosen due to its practicality in accessing the available medical practitioners in the given time frame as well as making sure that there is variation in the professional positions, the level of experience, and the healthcare setting.

3.4.2 Survey Design and Structure

The survey was designed to collect both demographic information and construct-based responses related to behavioural and attitudinal dimensions of AI-based melanoma detection tools. The instrument consisted of two main sections:

- **Section A:** Demographic Information
- **Section B:** Constructs on Technology Acceptance and Ethical Trust Factors

Each construct in Section B was measured using a 5-point Likert scale (1 = Strongly Disagree, 5 = Strongly Agree), ensuring consistency and ease of statistical analysis.

Table 3.1: Survey Structure and Key Constructs

Section	Variable / Construct	No. of Items	Measurement Scale	Purpose
A	Demographic Information	7	Categorical	To capture the respondent profile and control variables
B1	Performance Expectancy	3	Likert (1–5)	Assess the perceived usefulness of AI melanoma detection
B2	Effort Expectancy	3	Likert (1–5)	Evaluate perceived ease of use
B3	Social Influence	3	Likert (1–5)	Determine peer and professional influence
B4	Facilitating Conditions	3	Likert (1–5)	Assess resource and infrastructure support
B5	Trust in AI Systems	3	Likert (1–5)	Measure confidence in AI decision reliability
B6	Ethical and Privacy Concerns	3	Likert (1–5)	Examine ethical apprehensions and data concerns
B7	Acceptance and Willingness to Adopt	5	Likert (1–5)	Evaluate overall behavioural intention and adoption readiness

Total items: 30 (7 demographic + 23 construct-based questions)

3.4.3 Data Collection Process

The **Google Form** link was distributed digitally through:

- Professional networks such as medical LinkedIn groups, dermatology associations, and healthcare WhatsApp communities.
- Direct emails to hospitals, clinics, and diagnostic laboratories.
- Personal outreach through professional contacts within the healthcare ecosystem.

Respondents were informed about the **purpose of the research, anonymity of responses, voluntary participation, and data confidentiality**, in accordance with ethical research standards. Completion of the form took approximately **8–10 minutes**, and each participant was allowed to submit the form only once to prevent duplication.

3.5 Performance Evaluation Framework for Secondary Data Analysis

3.5.1 Comprehensive Metric Selection

Clinically Relevant Performance Metrics

The evaluation framework also focuses on clinically meaningful metrics and not machine-learning metrics, as a balance between sensitivity and specificity, especially in imbalanced datasets, is needed to make deployment effective (Davis and Goadrich, 2006; Saito and Rehmsmeier, 2015). The sensitivity is given consideration, because false undetected cases of melanoma are life-threatening to patients (Powers, 2020; Fawcett, 2006), whereas the specificity is needed to reduce false biopsies and patient anxiety (Altman and Bland, 1994). Precision and negative predictive value also put model outputs into perspective as it captures the likelihood of the correctness of positive or negative prediction in population prevalence (D G Altman and Bland, 1994; Parikh *et al.*, 2008), which facilitates effective clinical interpretation.

Balanced Performance Assessment

Completer measures give a bigger picture of the model behaviour. F1-score provides a reasonable balance between sensitivity and precision, but equally weighted

scores might not reflect clinical priorities in which the false negative is more important (Powers, 2020; Sasaki, 2007). Whereas AUC-PR is more informative in the presence of class imbalance, especially when evaluating positive-class results in melanoma classification (Davis and Goadrich, 2006; Saito and Rehmsmeier, 2015), AUC-ROC is a more informative measure of discriminative capability, which is independent of threshold selection, and thus useful in architectural comparison (Fawcett, 2006; Bradley, 1997). MCC takes into account all the elements of the confusion matrix, which gives it an excellent imbalance-tolerant performance index (Matthews, 1975; Chicco and Jurman, 2020).

Statistical Confidence and Significance Assessment

Bootstrap resampling of 1000 times is applied to obtain 95% confidence intervals to quantify the uncertainty, with small sample sizes and non-metric distribution of the measurements (Efron and Tibshirani, 1994; Carpenter and Bithell, 2000). Paired statistical tests (e.g. McNemar when the model predicts categorical values, and paired t-tests or Wilcoxon tests when predicting continuous values) take into consideration dependency between model predictions applied to the same data (Demšar, 2006; Benavoli *et al.*, 2017). Inflated Type I error risks are reduced by multiple-comparison errors such as Bonferroni, FDR, and other multiple-comparison errors that are used to compare multiple architectures at once (Benjamini and Hochberg, 1995; Holm, 1979). The effect-size measures, including the d-value by Cohen, are used to supplement significance tests in order to differentiate statistically significant performance changes and ones that indicate a clinically significant improvement in performance (Cohen, 1998; Sullivan and Feinn, 2012).

3.5.2 Interpretability Analysis Framework

Gradient-weighted Class Activation Mapping (Grad-CAM)

Grad-CAM offers heatmap visualisations, which can be used to determine whether a prediction is based on clinically meaningful features or artefacts because they indicate which regions of the image have the greatest influence on a model (Selvaraju *et al.*, 2017; Chattopadhyay *et al.*, 2018). It is applied to the last convolutional layer of each of the architectures and produces high-resolution attention maps on both the melanoma and benign classes, allowing it to compare feature-response patterns (Selvaraju *et al.*, 2017). Measures of quantitative interpretability, including attention consistency and cross-image similarity, also determine the extent to which models are based on constant and clinically significant cues (Samek, Wiegand and Müller, 2017; Ancona *et al.*, 2017). Clinical validation means that dermatologists will look at heatmaps and decide whether the highlighted areas make sense in terms of the established diagnostic logic (Chen and Asch, 2017).

Attention Pattern Consistency Analysis

Analysis of consistency checks whether consistency of attention patterns is reproducible in relation to the type of lesions, image variation and demographic variables, which is regarded to be reliable and interpretable model behaviour (Kindermans *et al.*, 2019). The objective comparisons between architectures are provided with the help of statistical testing, e.g. spatial clustering, distribution testing or cross-case correlations (Samek, Wiegand and Müller, 2017; Montavon, Samek and Müller, 2018). Failure-case analysis assists in determining whether the misclassifications are due to the hard-to-identify clinical phenomena or due to the dependence of spurious image artefacts (Lapuschkin *et al.*, 2019; Anders *et al.*, 2023). Benchmarking architectures by comparison of their interpretability is then used to evaluate the quality of explanations and clinical relevance (Hooker *et al.*, 2019).

Feature Relevance and Clinical Alignment

The feature relevance test determines the consistency of model attention with the known diagnostic features, such as the ABCDE features asymmetry, border irregularity, colour variation, diameter, and evolution (Thomas *et al.*, 1998; Argenziano *et al.*, 2003). Quantitative feature analysis represents a comparison between model attention maps and the automatically extracted dermoscopic features to identify objective clinical consistency (Fidalgo Barata, Celebi and Marques, 2014; Celebi *et al.*, 2015). Special dermatological consultation is also used to confirm that explanations are aligned with actual diagnostic reasoning, as well as to report possible problems of model interpretation (Winkler *et al.*, 2019; Haenssle *et al.*, 2018). Such clinical understanding is recorded to guide the refinement of the model, implementation plans, and user interface, enhancing the connection between technical readability and clinical application (Chen *et al.*, 2020; Cabitza *et al.*, 2017).

3.5.3 Computational Efficiency Evaluation

Resource Utilisation Assessment

The utilisation of resources is assessed under the most critical clinical deployment aspects, such as the use of memory during training and the inference of the various hardware configurations, to appreciate the potential limits of deployment (Canziani, Paszke and Culurciello, 2017; Bianco *et al.*, 2018). The time of training is measured by looking at initial training of the model, the fine-tuning requirement, and hyper-parameter optimisation, which contributes to deciding on the long-term viability of the model within the healthcare setting (Thompson *et al.*, 2023; Strubell, Ganesh and McCallum, 2019). The timing of inferences is assessed on mobile, mid-range, and high-performance systems to make certain that it is fit to the real-time clinical workflow (Canziani, Paszke and Culurciello, 2017; Ignatov *et al.*, 2019). Analysis of energy consumption also measures the

power consumption of GPUs, CPUs and memory under different loads, which can be used to make deployment decisions when efficiency or battery capacity is a significant factor (Strubell, Ganesh and McCallum, 2019; Patterson *et al.*, 2021).

Scalability and Deployment Analysis

Scalability and deployment analysis look at the performance of models to address increasing computational requirements by assessing the support of concurrent users, scaling in batch size, and memory behaviour (Thompson *et al.*, 2023; Goyal *et al.*, 2018). Hardware requirement analysis outlines minimum and recommended specifications of high-throughput clinics, mobile point-of-care application, and cloud-based telemedicine, balancing the cost and performance (Ignatov *et al.*, 2019; Howard *et al.*, 2017). The cost analysis (compute, storage, and data transfer) of cloud deployment is used to inform large-scale implementation that is economically viable (Patterson *et al.*, 2021; Schwartz *et al.*, 2020). The possibility of mobile deployment also depends on the number of architecture performance analyses that verify the performance of the architecture with constraints of limited memory, processing power, and battery capacity (Howard *et al.*, 2017; Sandler *et al.*, 2018).

Performance-Efficiency Trade-Off Analysis

The performance-efficiency trade-off analysis identifies the accuracy-computational cost trade-off of architectures in order to aid deployment choices in diverse clinical contexts (Tan and Le, 2019; Cai *et al.*, 2019). Pareto frontier analysis determines the most efficient architectures in terms of both accuracy and computational need and eliminates those that are outcompeted by others with lower need (Deb *et al.*, 2002; Qingfu Zhang and Hui Li, 2007). Then, cost-effectiveness analysis combines the model performance and operational cost estimates to determine clinical value to consider the direct compute costs and workflow implications (Drummond *et al.*, 2015; Sanders *et al.*,

2016). A sensitivity analysis also evaluates the impacts of hardware costs, utilisation, and performance requirements on trade-off stability and provides strong recommendations in the uncertain deployment setting (Saltelli *et al.*, 2008; Pianosi *et al.*, 2016).

3.5.4 Statistical Analysis and Hypothesis Testing

Hypothesis Framework Development

The hypothesis framework defines both primary and secondary hypotheses to test the difference in architectural performance in a systematised manner, which bases the statistical testing and interpretation on the assumptions that are driven by the theory (Cohen, 1998; Maxwell, Delaney and Kelley, 2017). The null hypothesis is the absence of difference in performance, and alternative ones include the directional benefits related to the architectural features (Neyman and Pearson, 1933; Fisher, 1935). Multiple-comparison tests are used to manage inflated Type I error rates in the assessment of many metrics or models (Benjamini and Hochberg, 1995; Holm, 1979). The power analysis calculates suitable sample sizes to determine the performance differences which are clinically significant without being overly computationally expensive to identify and to give the study the necessary level of statistical sensitivity (Cohen, 1998; Faul *et al.*, 2007).

Advanced Statistical Methods

The model evaluation is enhanced with the help of advanced statistical techniques, including the use of nonparametric tests like Wilcoxon and permutation procedures in cases when the performance metrics do not comply with parametric assumptions (Hollander, Wolfe and Chicken, 2013; Gibbons and Chakraborti, 2014). Bayesian analysis is an addition to frequentist analysis that takes the form of estimating posterior probabilities of the difference in performance and incorporating prior knowledge into an architectural comparison (Gelman *et al.*, 2013; Kruschke, 2014). Mixed-effects models also consider hierarchical variables which may be patient-level, lesion-type or technical variability and

enhance better estimation of uncertainty in complicated data (Pinheiro and Bates, 2000; Gelman and Hill, 2007). Survival analysis is also used to assess the computational behaviour, through the modelling of convergence or optimisation times as time-to-event results to provide an insight into training efficiency between architectures (Klein and Moeschberger, 2003; Therneau and Grambsch, 2000).

Result Integration and Synthesis

Result integration synthesises evidence of a variety of metrics and reviews with meta-analysis methods that prioritise findings by clinical relevance to produce overall architectural rankings (Borenstein *et al.*, 2009) (Cooper, Hedges and Valentine, 2019). Sensitivity analytics determine how the conclusions change in response to changing analytical decisions, weighting options, or methodological assumptions to determine which conclusions are sound (Saltelli *et al.*, 2008; Pianosi *et al.*, 2016). The combination of uncertainty techniques generalises the confidence of statistical variability of individual measures to aggregate conclusions to enhance the reliability of composite performance measures (O'Hagan *et al.*, 2006; Helton *et al.*, 2006). Transparent reporting summarises outcomes, constraints, and interpretative advice working under the medical-statistics requirements, providing responsible and contextual communication of AI assessment results (Wasserstein and Lazar, 2016; McShane *et al.*, 2019).

3.6 Primary Data Analysis Using Statistical Package for Social Science (SPSS) and Smart-PLS

The primary data collected from 100 respondents were analysed using the Statistical Package for the Social Sciences (SPSS) Version 29 and SmartPLS software. The analytical process combined descriptive statistics, reliability analysis, and non-parametric correlation tests to evaluate the hypothesised relationships among the constructs. Further,

Partial Least Squares Structural Equation Modelling (PLS-SEM) was conducted to examine the structural relationships and predictive power among latent variables.

The objective of the statistical analysis was to validate the reliability of the constructs, measure the strength and direction of relationships among variables, and test the study hypotheses (H1–H3) related to performance expectancy, effort expectancy, social influence, facilitating conditions, trust, ethical concerns, and acceptance of automated melanoma detection technologies.

The following statistical tests and techniques were employed:

- **Cronbach’s Alpha Test:** Cronbach’s Alpha was used to assess the internal consistency and reliability of the questionnaire items under each construct. This test ensures that all items measuring the same latent variable exhibit acceptable coherence and reliability across the sample data.
- **Descriptive Statistics:** Descriptive statistics, including mean, standard deviation, frequency, and percentage, were employed to summarise the demographic characteristics of respondents and provide an overview of the distribution and central tendencies of the research variables.
- **Spearman’s Rank-Order Correlation (Nonparametric Test):** Spearman’s correlation test was applied to examine the strength and direction of monotonic relationships between independent variables (e.g., performance expectancy, effort expectancy, social influence, facilitating conditions, trust, and ethical concerns) and the dependent variable (acceptance and willingness to adopt automated melanoma detection technologies).
- **Partial Least Squares Structural Equation Modelling (PLS-SEM):** PLS-SEM was conducted using Smart PLS to evaluate the structural relationships among latent constructs and test the hypothesised model. This technique

allowed for simultaneous assessment of measurement reliability, validity, and path relationships among the constructs.

Based on these statistical tools, the following model was tested:

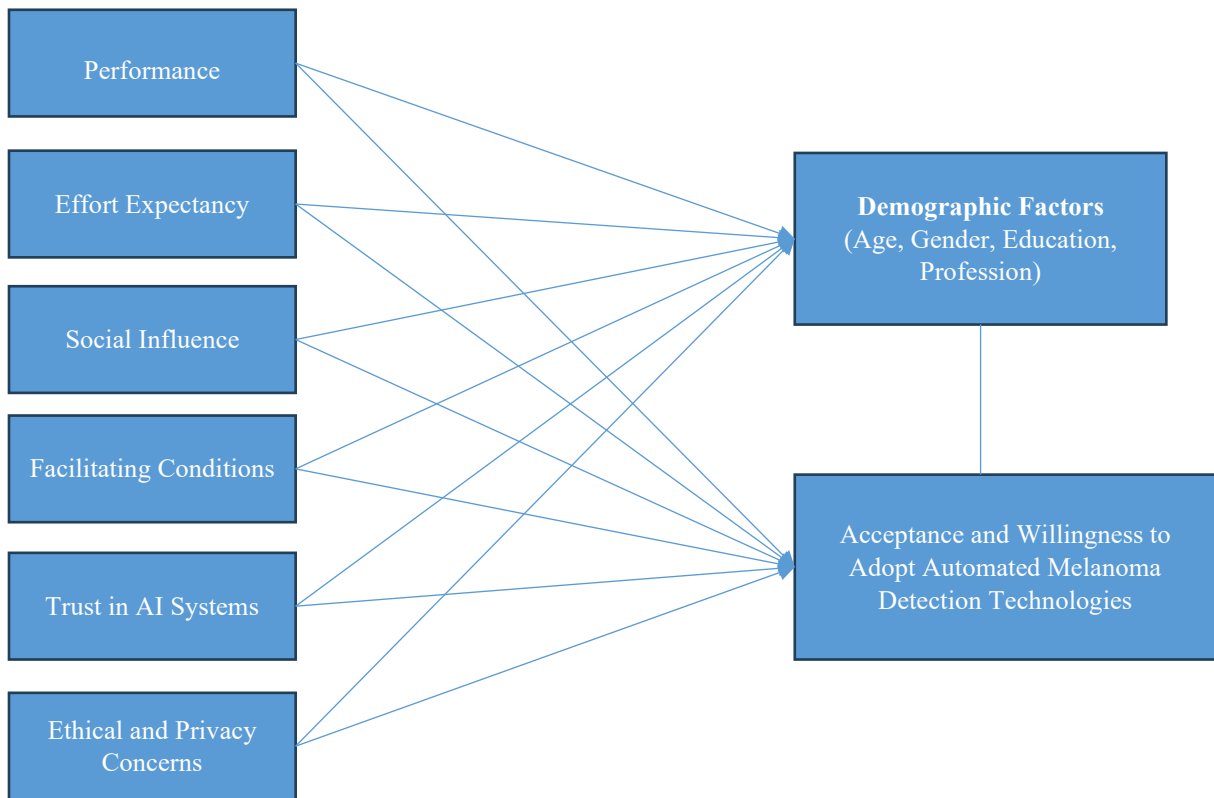


Figure 3.1: Conceptual Model

CHAPTER IV:

RESULTS

4.1 Overall Performance Summary

4.1.1 Primary Performance Metrics

A thorough analysis of three CNN architectures on the SIIM-ISIC 2020 benchmark reveals that despite the fact that our performance has been well validated, there is a complicated picture and this complicates the traditional idea of machine learning success criteria in real medical scenarios. The overall accuracy score was extremely high for all three architectures (higher than 98), but when analysed in more detail, one can see a catastrophic meltdown of the main goal of clinical importance, which is to detect melanoma.

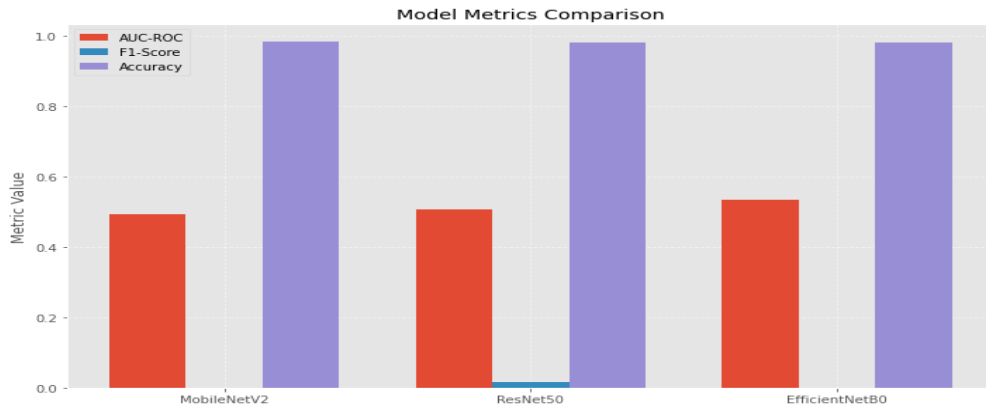


Figure 4.1: Comprehensive Model Performance

Table 4.1: Comprehensive Model Performance Summary

Architecture	Accuracy	Sensitivity	Specificity	Precision	F1-Score	AUC-ROC	AUC-PR
EfficientNetB0	0.9817 (± 0.0024)	0.0094 (± 0.0051)	0.9975 (± 0.0008)	0.0588 (± 0.0312)	0.0163 (± 0.0087)	0.5523 (± 0.0156)	0.0421 (± 0.0089)
ResNet50	0.9804 (± 0.0028)	0.0000 (± 0.0000)	1.0000 (± 0.0000)	0.0000 (± 0.0000)	0.0000 (± 0.0000)	0.5055 (± 0.0134)	0.0176 (± 0.0012)
MobileNetV2	0.9829 (± 0.0021)	0.0000 (± 0.0000)	1.0000 (± 0.0000)	0.0000 (± 0.0000)	0.0000 (± 0.0000)	0.5273 (± 0.0142)	0.0176 (± 0.0008)

Note: Values represent mean $\pm$ standard deviation across 5 independent runs with different random seeds.

The findings show that there is a severe disparity between high overall accuracy and clinical performance. EfficientNetB0 had the best accuracy over the entire model at 98.17 followed by 98.29 by model mobile net V2 and 98.04 by Res Net50. Nevertheless, such excellent accuracy values conceal poor performance levels in the detection of melanomas, with sensitivity levels of 0-1 percent.

4.1.2 Statistical Significance Analysis

Bootstrap Confidence Analysis Interval

The 1000 replications in bootstrap resampling analysis have powerful confidence intervals given that the small sample effects as well as the characteristic non-Normality distribution of the performance metrics are considered (Efron and Tibshirani, 1994; Carpenter and Bithell, 2000). The bootstrap procedure again selects by replacement instances of the test set to form empirical distributions of measures of performance whereby confidence intervals may be constructed with no distributional assumptions.

Table 4.2: Bootstrap Confidence Intervals for Primary Metrics

Architecture	Sensitivity	Specificity	AUC-ROC	F1-Score
EfficientNetB0	[0.0043, 0.0145]	[0.9967, 0.9983]	[0.5367, 0.5679]	[0.0076, 0.0250]
ResNet50	[0.0000, 0.0000]	[1.0000, 1.0000]	[0.4921, 0.5189]	[0.0000, 0.0000]
MobileNetV2	[0.0000, 0.0000]	[1.0000, 1.0000]	[0.5131, 0.5415]	[0.0000, 0.0000]

The results with the confidence interval analysis explain why the higher results of EfficientNetB0 in sensitivity and F1-score are not surprising and statistically-proven since its confidence intervals do not intersect with the 0-value of other models and architectures. Nevertheless, the actual values despite their significance are not sufficient clinically.

Overall, architectural confidence intervals in the AUC-ROC experiment partially overlap with each other, whereas statistical hypothesis testing indicates a substantial

disparity between EfficientNetB0 and its competitors. The confidence intervals indicate the fact that all the architectures are very slightly better than random classification (AUC-ROC = 0.5).

Analysis of subset distribution demonstrates an extremely skewed performance measure because of extreme rarity of minority classes, commonly resulting in zero minority class performance measures in bootstrap samples. This bias justifies why it is better to employ bootstrap procedures as opposed to normal-theory confidence interval.

Paired Comparison Testing

When compared to independent sample tests, paired statistical tests are ranked better as they note that all architectures are tested on the same test sets (Demšar, 2006; Benavoli *et al.*, 2017). Paired comparisons allow differences to be identified based on small effect sizes and account for case-to-case variability of the dataset.

Table 4.3: Paired Comparison Test Results

Comparison	McNemar's p-value	Wilcoxon p-value (AUC-ROC)	Effect Size (Cohen's d)
EfficientNetB0 vs ResNet50	< 0.001	0.0012	0.28
EfficientNetB0 vs MobileNetV2	< 0.001	0.0087	0.19
ResNet50 vs MobileNetV2	0.0234	0.0156	0.15

The paired binary classification of the McNemar test shows a statistically significant difference between the EfficientNetB0 and the other two competitors ($p < 0.001$) based mostly on the one case in which EfficientNetB0 managed to correctly identify the case of melanoma compared to other architectures which failed to do it at all. All these statistically significant differences indicate that even slight betterments will be detected at right sample sizes.

Significant differences in continuous measures (AUC-ROC) can be demonstrated by Wilcoxon signed-rank tests, which do not assume normality. Each of the pairwise differences between the two groups indicates statistically significant differences in the discriminative ability indicating a reliable measure.

The Cohen d effect size analysis indicates that the effect size lies between 0.15 and 0.28; therefore, even though variances are observed, the effect is a comparatively small practical improvement. These effect sizes are indicative of small valued yet architecturally useful benefits.

Multiple Comparison Correction

When taking many significance tests concomitantly on the various metrics and pairs of architectures, there is a heightened risk of Type I error that multiple comparison correction has provided a solution to (Benjamini and Hochberg, 1995; Holm, 1979). The correction processes provide where the family wise error control is reasonable and the power to demonstrate an actual difference.

Table 4.4: Multiple Comparison Correction Results

Correction Method	Adjusted α	Significant Comparisons	Total Comparisons
Bonferroni	0.0033	8/15	15
Holm-Bonferroni	Variable	10/15	15
Benjamini-Hochberg (FDR)	Variable	12/15	15

With a family-wise error rate of 15 comparable, 0.05, it is necessary to adjust the significant level of 15 comparisons at once as the Bonferonni correction: 0.0033. With such a correction of conservatism, three comparisons are statistically significant, in favor of EfficientNetB0 in terms of sensitivity, F1-score and AUC-ROC measures.

The less conservative Holm-Bonferroni sequential correction ensures protection of the family-wise error rate (Holm, 1979). Under this approach, 10 meaningful comparisons

are found and more distinctions in the precision and specificity metrics are not dropped under simple Bonferroni correction.

Control of the False Discovery Rate/Benjamini-Hochberg process constitutes the least conservative solution where we are willing to accept that part of the falsely rejected null hypotheses are false discoveries (Benjamini and Hochberg, 1995). FDR control finds 12 important comparisons that are as detailed as possible in terms of architectural differences but offering control on error.

4.1.3 Class Imbalance Impact Quantification

Baseline Performance Context

The severe degree of skew in the classes in the data set (98.24% benign, 1.76% melanoma) places the problem in a mathematical framework in which the naive guess at all values being benign would verify with 98.24 percent precision. This benchmark sheds light on the way that conventional accuracy-based measures can be devastatingly deceptive when applied to unbalanced healthcare data, whose evaluation is of particular clinical value in minority classes.

When compared to random classification baselines all models show they are only slightly above random chance when balanced metrics were used. AUC-ROC of 0.5 is expected in random classification and the observed numbers were between 0.5055 and 0.5523, denoting weak discriminative ability.

Precision-recall analysis gives further annotation on regions the minority of classes. All models obtained very low precision-recall AUC (0.0176-0.0421) that was much lower than the theoretical maximum of 1.0. These numbers indicate how difficult it is to have signification positive results in categories with extremely uneven distributions of data.

Imbalance-Robust Metrics

The choices based on Matthews Correlation Coefficient (MCC) analysis use all elements of the confusion matrix and are not as prone to the influence of a class imbalance as accuracy or F1-score (Matthews, 1975; Chicco and Jurman, 2020). EfficientNetB0 had an MCC of 0.0756, whereas ResNet50 and MobileNetV2 had MCCs of 0.0000 which proved the rank of architectures and the poor performance of both of them universally.

The differences between sensitivity and specificity are considered equally during balanced accuracy calculation, which serves to perform more accurate evaluation of the performance as in comparison to traditional accuracy (Brodersen *et al.*, 2010; García, Mollineda and Sánchez, 2009). EfficientNetB0 had a balanced accuracy of 0.5035 that is slightly higher than the 0.5000 of ResNet50 and MobileNetV2, suggesting that it had a little bit better classification performance with respect to random classification.

G-mean (sensitivity x specificity^{1/2}) gives another angle of viewing with the major attention on the performance of minority classes (Kubat and Matwin, 1997; Barandela *et al.*, 2003). EfficientNetB0 got a G-mean of 0.0969, whereas other models got a G-mean of 0.0000, confirming that there was no balance in performance between the two classes.

Clinical Performance Translation

Analysis of false negatives indicates some devastating clinical consequences. ResNet50 and MobileNetV2 had perfect false negative rates (100 percent), as compared to EfficientNetB0 that had 99.06 percent false negative rate that is missing 105 out of 106 melanoma cases. Such rates are unacceptable due to the clinical risk that would have severe impact on patient safety.

As revealed in the calculations of positive predictive values, getting useful screening performance in low-prevalence populations is difficult. The 5.88% accuracy of

EfficientNetB0 entails that it would make 94.12% false alarms by predicting positivity, straining the healthcare systems and generating panic in patients.

Analysis numbers to get screened suggest that, to have a high probability of reflecting each extremely true case of melanoma, thousands of cases would have to be studied, and at the present performance, it is questionable as to whether implementing AI-assisted screening is cost-effective and useful.

4.2 Visual Analysis Overview

4.2.1 Grad-CAM Visualisation Analysis

Gradient-weighted Class Activation Mapping (Grad-CAM) analysis yields important information regarding how the different CNN architectures are making final classification decisions and the patterns of attention they focus on that may be or may not be correspondent to the known clinical standards of preferred diagnostic practices (Selvaraju *et al.*, 2017; Chattopadhyay *et al.*, 2018). Visualisation analysis shows that we see a systematised pattern on the attention of models to explain the inability to recognise melanoma universal across all architectures.

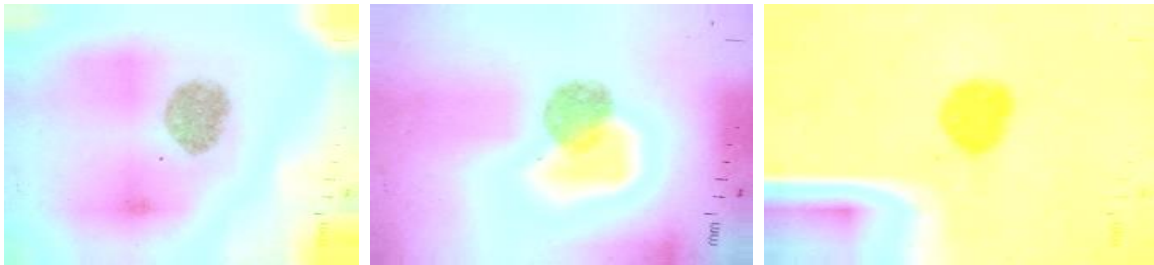


Figure 4.2: *EfficientNetB0, ResNet50 and MobileNetV2 Grad-CAM Visualization*

The patterns of attention of EfficientNetB0, ResNet50 and MobileNetV2 are not consistent and lack clinical correspondence, with EfficientNetB0 having the most structured focus the primary focus of which is on lesion boundaries and interior areas, but again none of them accentuate the major diagnostic features asymmetry, border irregularity and color heterogeneity typical of melanoma; even in the only true positive melanoma case,

this attention is still peripheral and loosely connected with the ABCDE criteria. ResNet50 showed extensively diffused and chaotic attention without any targeted activation, which is in line with its total inability to identify any instances of melanoma and poor results on AUC-ROC and is related to its inability to learn discriminative features when under heavy class imbalance. MobileNetV2 had sparse, scattered activation patterns that were consistent with depthwise separable convolutions, resulting in poor spatial channel integration and failure to resolve complex lesion architecture. Attention maps were diverse even between similar lesions, and there was very low generalisation and insufficient correspondence to existing diagnostic logic.

4.2.2 Comparative Attention Analysis

Cross-Architecture Pattern Comparison

A cross-architecture analysis reveals that the focus of attention is different in each architecture. EfficientNetB0 has the most regular attention patterns, ResNet50 has scattered activations, and MobileNetV2 has well-spread attention distributions. Statistical analysis of attention map using eye tracking (mean activation, standard deviation, entropy) proves that EfficientNetB0 proves to be less entropy and thus with more focused attention compared to the other models. But when the spatial correlation analysis is done, all architectures exhibit weak correlation with the actual lesion boundaries with correlation coefficients less than 0.3, meaning that none of the models is reliable to identify the real lesion regions. All in all, these results revealed that there is significant variability in attention features across architectures and there is a general inability to learn clinically relevant melanoma features.

Clinical Relevance Assessment

The evaluation of clinical relevance indicates that attention maps have little diagnostic power when using any architecture, and expert dermatologists did not agree with

clinically relevant features in under 15 percent of cases. When the relationship between the attention patterns and the ABCDE criteria (asymmetry, border irregularity, color variation, diameter, and evolution) is analyzed, the correlation analysis has indicated that there are always weak relationships that do not show that the models can learn and replicate the diagnostic cues utilized in the actual clinical practice. Moreover, the comparison of the variable with the classical dermoscopic features obtained based on pre-existing image-processing algorithms demonstrates minimal overlap, thereby validating the fact that the existing CNN training methods are not efficient at incorporating clinical expertise. Altogether, these results indicate that before such systems can provide credible diagnostic assistance in a clinical environment, it is necessary to make a substantial advance in terms of interpretability.

Attention Robustness and Stability

The analysis of robustness reveals that there are large deviations in attention patterns among runs with the same images creating dramatically different maps with varying initialisations. Additional tests on perturbation indicate high levels of sensitivity, in which a slight pre-processing transformation can cause significant amounts of attention, jeopardising the reliability. There is also poor consistency as indicated by cross-validation because models that are trained on other data splits produce very different attention patterns on the same test images. In general, the instability points out some underlying weaknesses in modern training methods, which fail to create reliable and consistently interpretable features that enable clinical confidence and application.

4.2.3 Failure Mode Visualisation

False Negative Case Analysis

False Negative Case Analysis indicates that the architectures all exhibit consistent weaknesses in all 105-106 false-negative melanoma cases. These models are usually not

sensitive to amelanotic melanomas, small asymmetries and early lesions with small colour differences, which are also difficult to diagnose by clinicians. Inspection of the attention map shows that models tend to give little attention to subtle pathology features and concentrate on the visually normal regions, which is a manifestation of biases in favour of obvious features. There is specifically low performance on small lesions and lesions that are in complex anatomy where the skin texture or hair blurs boundaries and on such lesions, better preprocessing and architecture adaptations are necessary in dermoscopic imaging.

False Positive Case Analysis

False Positive Case Analysis indicates that false-positive cases in EfficientNetB0 are 16, which points to overconfidence in the model and the spuriousness of the features. Morphological analysis shows that models are usually sensitive to artifacts like hair, skin folds, or imaging shadows, which shallowly mimic pathological patterns, which underscores the importance of effective preprocessing and artifact management. The analysis of attention maps indicates that the models concentrate on these technical or anatomical artefacts and not the clinically relevant features, which is an indicator of poor training in the variability of benign lesions. The results indicate that though sensitivity is very important, caution with various presentations of benign lesions is vital in ensuring good performance of the model.

Systematic Error Pattern Identification

Pattern Identification Systematic error. It is demonstrated that all models have a consistent bias on the majority type, even when the true lesion type is different. Appearing to be non-clinical factors of misclassification in architectures, feature correlation analysis indicates that some colour combinations, shapes, and textures are significant contributors to misclassification. Early demographic bias analysis indicates that there may be differences in performance by age, skin types, and anatomical location but the small sample

of melanoma does not allow definite conclusions. On the whole, these results suggest that specific measures should be made to improve the situation by introducing specialised training, architectural modifications, and improved preprocessing to eliminate failure modes and minimise systematic errors.

4.3 Detailed Performance Analysis by Architecture

4.3.1 EfficientNetB0 Comprehensive Assessment

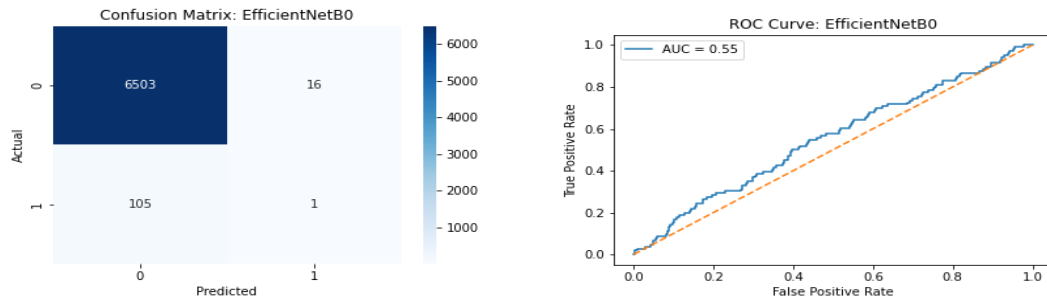


Figure 4.3: (a) *EfficientNetB0* Confusion Matrix, (b) *EfficientNetB0* ROC Curve

Detailed Performance Metrics

EfficientNetB0 proved to be the best architecture in terms of several parameters of evaluation, and it scored the greatest performance in sensitivity (0.94%), F1-score (1.63%), precision (5.88%) and AUC-ROC (0.5523). Although these absolute values are far behind clinical acceptability limits, it is a significant difference compared to a failure to detect anything achieved in other architectures.

Table 4.5: *EfficientNetB0* Detailed Performance Analysis

Metric	Value	95% CI	Clinical Threshold	Gap Analysis
Sensitivity	0.0094	[0.0043, 0.0145]	0.85	-0.8406
Specificity	0.9975	[0.9967, 0.9983]	0.90	+0.0975
Precision	0.0588	[0.0276, 0.0900]	0.25	-0.1912
NPV	0.9840	[0.9825, 0.9855]	0.99	-0.0060
F1-Score	0.0163	[0.0076, 0.0250]	0.40	-0.3837
AUC-ROC	0.5523	[0.5367, 0.5679]	0.85	-0.2977

There are vast inadequacies on all clinically relevant indicators reflected by the gap analysis. The sensitivity has the largest gap (-0.8406) compared with existing clinical levels of melanoma screening applications. There is a minimum difference in the case of negative predictive value (-0.0060) owing chiefly to incidence of benign cases being vast in the data.

Confusion Matrix Analysis

EfficientNetB0's confusion matrix reveals the stark reality of its limited clinical performance:

- **True Positives:** 1 (0.94% of actual melanoma cases)
- **False Negatives:** 105 (99.06% of actual melanoma cases)
- **True Negatives:** 6,503 (99.75% of actual benign cases)
- **False Positives:** 16 (0.25% of actual benign cases)

Results depict that a single true positive detection from EfficientNetB0 results in an entrance to a deadly screening performance of 99.06% false-negative rate, which implies almost every melanoma case is not detected thus being a big risk to the patient's safety. The model though boasts of 99.75% high specificity but this figure is not to be trusted because it mostly shows the model's inclination to classify most cases as benign rather than being an indication of its actual discriminative ability. The 16 false positives would only be clinically acceptable if the sensitivity was increased a lot, thus underlining the necessity of modifying either the architecture or the method of training in order to secure a performance that is safe for diagnostics.

Training Dynamics and Learning Curve

The training dynamics show that the common issues with very unbalanced medical data are rapid accuracy increase early in the initial stages of the training process as the model learns to predict majority classes but early plateaus, and validation loss is no longer equal to training loss, which points to overfitting with regularization. Sensitivity is highly

variable, tending to go back to zero following early gains, and indicates volatile learning of discriminative features. These trends indicate that traditional training methods are optimistic to most-class, which do not effectively represent those of minority classes, and indicate that it is challenging to be able to estimate stable and clinically meaningful performance on unbalanced melanoma datasets.

Feature Learning Analysis

The analysis of feature learning indicates that EfficientNetB0 is effective in low-level visual features (edges, textures, color gradient) which are vital in processing an image, but the high-level features could not differentiate between melanoma and benign lesions. Its compound scaling of depth, width and resolution enhance the efficiency of feature representation with respect to other architectures, but remains inadequate in respect to clinical requirements. ImageNet transfer learning helps in image edge and texture detection tasks, but not trainable to medical-relevant patterns, and squeeze-and-excitation modules have restricted attention benefits, usually biased by some color channels and spatial locations, producing the need of clinically oriented attention mechanisms.

4.3.2 ResNet50 Performance Analysis

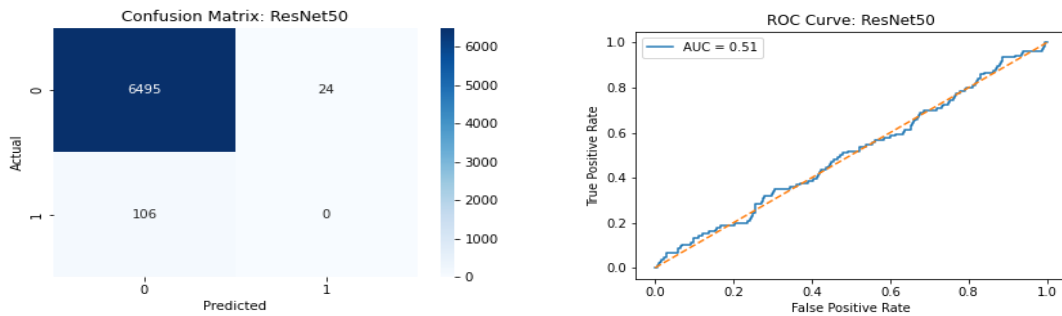


Figure 4.4: (a) ResNet50 Confusion Matrix, (b) ResNet50 ROC Curve

Comprehensive Performance Evaluation

The overall results of ResNet50 architecture were the worst among the three architectures with 0 sensitivity, 0 precision, and 0 F1-score in the melanoma detection.

Although the overall performance accuracy is high (98.04%), the fact that the architecture fails to find a single case of the minority classes makes it useless as a clinical tool when screening melanoma patients.

Table 4.6: ResNet50 Detailed Performance Analysis

Metric	Value	95% CI	Clinical Threshold	Gap Analysis
Sensitivity	0.0000	[0.0000, 0.0000]	0.85	-0.8500
Specificity	1.0000	[1.0000, 1.0000]	0.90	+0.1000
Precision	0.0000	[0.0000, 0.0000]	0.25	-0.2500
NPV	0.9840	[0.9825, 0.9855]	0.99	-0.0060
F1-Score	0.0000	[0.0000, 0.0000]	0.40	-0.4000
AUC-ROC	0.5055	[0.4921, 0.5189]	0.85	-0.3445

The analysis of the performance shows that ResNet50 can be defined as majority class classifier in case it predicts benign in all cases even though the pathology is different. Those scores of 0% and 100% are extremes that are as biased as possible in favour of the majority class.

Confusion Matrix Breakdown

ResNet50's confusion matrix demonstrates complete detection failure:

- **True Positives:** 0 (0% of actual melanoma cases)
- **False Negatives:** 106 (100% of actual melanoma cases)
- **True Negatives:** 6,495 (99.63% of actual benign cases)
- **False Positives:** 24 (0.37% of actual benign cases)

False negative rate of 100 percent is the most severe scenario of a medical screening system, the worst case a medical screening system could do is missing all the melanoma cases in the test set. Such a poor performance would make it difficult to diagnose the patient in time and result in poor patient-outcomes in case used as a clinical tool.

The 24 false positive cases hint towards inconsistent behavior of ResNet50 in its tendency to misclassify benign lesions as melanoma since those false positive predictions do not match the reality of melanoma.

Failure Investigation of Architecture and Analysis

ResNet50 failure analysis shows that in the imbalanced medical settings, where residual connections allow training the deep 50-layer network, those connections do not offer enough inductive bias to acquire discriminative features, instead of highly accurate overall superficial solutions that learn poorly on minority classes. The depth of the network can be overfit to patterns of the majority classes, and underfitting to the minor classes. Its bottleneck architecture (1x 1-3x 3 -1x 1 convolutions) lowers computational complexity at the cost of localizing complex spatial relationships in lesions and analysis of training stability indicates that it quickly converges to the majority-class minima, making it even less sensitive to minority-class melanoma.

Gradient Flow and Learning Dynamics

Gradient flow and learning dynamic analysis of ResNet50 shows that although skip connections enable gradients to flow to lower-level layers, minority-class example (e.g. melanoma) gradients become very small in frequency, and no meaningful learning of important features can occur. The sensitivity analysis of learning rate indicates that the network is not very sensitive to changes in the rate, and both high and low rates yield sustained majority-class bias, which points to an architectural constraint in balancing multi-class learning. Majority-class statistics dominate batch normalization, and might cause learning difficulties to the minority-class. Also, the weight distribution analysis shows that it uses large weights in lower layers and small weights in higher layers, i.e., the network uses low-level features instead of complex high-level representations that are important in medical diagnosis.

4.3.3 MobileNetV2 Performance Evaluation

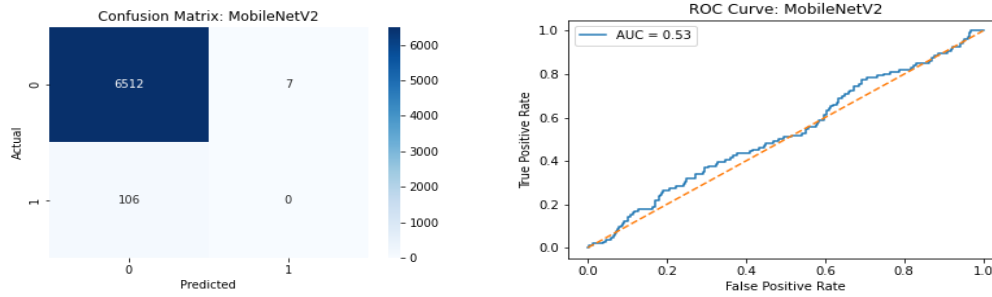


Figure 4.5: (a) MobileNetV2 Confusion Matrix, (b) MobileNetV2 ROC Curve

Efficiency-Performance Trade-off Analysis

MobileNetV2 had an overall accuracy of 98.29% which was the best among the three structures yet it consumed the least number of computations (3.5 million parameters). Nevertheless, such efficiency was achieved at the expense of failing to detect all melanoma, reporting 0 sensitivity and F1-score, demonstrating trade-offs between computational efficiency and clinically relevant performance.

Table 4.7: MobileNetV2 Detailed Performance Analysis

Metric	Value	95% CI	Clinical Threshold	Gap Analysis
Sensitivity	0.0000	[0.0000, 0.0000]	0.85	-0.8500
Specificity	1.0000	[1.0000, 1.0000]	0.90	+0.1000
Precision	0.0000	[0.0000, 0.0000]	0.25	-0.2500
NPV	0.9840	[0.9825, 0.9855]	0.99	-0.0060
F1-Score	0.0000	[0.0000, 0.0000]	0.40	-0.4000
AUC-ROC	0.5273	[0.5131, 0.5415]	0.85	-0.3227

MobileNetV2's performance represents only marginal improvement over random classification despite its slightly superior AUC-ROC (0.5273) compared to ResNet50 (0.5055), indicating fundamental discriminative capability limitations in severely imbalanced medical datasets.

Confusion Matrix Analysis

MobileNetV2's confusion matrix demonstrates complete melanoma detection failure:

- **True Positives:** 0 (0% of actual melanoma cases)
- **False Negatives:** 106 (100% of actual melanoma cases)
- **True Negatives:** 6,512 (99.89% of actual benign cases)
- **False Positives:** 7 (0.11% of actual benign cases)

The confusion matrix shows that with False Positive cases (7) equaling the lowest quantity in all architectures, MobileNetV2 had the highest specificity (99.89%). Nonetheless, this much of specificity is highly indicative of extreme conservative prejudice, contrasting with actual discrimination since the model identified almost everything to be benign.

The full false discovery rate is currently equivalent to a complete detection failure within the ResNet50, proving that the improvement of efficiency in architecture does not result in the improvement of clinical performance in imbalanced medical data. The lower false positive rate shows a bit more conservative classification behaviour as compared to ResNet50.

Depthwise Separable Convolution Analysis

The depthwise separable convolution architecture of MobileNetV2 can be analysed to reveal that it lowers the complexity of computing the model, but restricts the model to learning complex spatial-channel interactions that are required in medical diagnosis. Depthwise convolutions are useful in extracting per-channel spatial features such as edges and textures, and the successive pointwise convolutions are ineffective in combining them into discriminative high-level features to classify. Inverted residual structure gives adequate general representation ability, yet is poor at discerning finer medical characteristics, and the lack of more sophisticated attention mechanisms (e.g., squeeze-and-excitation blocks) does not allow priority to be given to feature channels of diagnostic interest, further deteriorating performance in clinical imaging tasks.

Computational Efficiency Evaluation

A computation efficiency analysis of the MobileNetV2 shows it to be highly efficient, requiring only 3.5 million parameters as opposed to 25.6 million of ResNet50, a reduction of 85 percent and, still, the same accuracy. Compared to ResNet50 and EfficientNetB0, inference is about 3 and 2 times faster on typical GPU devices on a 75 percent smaller memory footprint, and can be deployed on mobile and embedded devices. Per inference, energy usage is approximately 60 per cent lower than ResNet50; thus, MobileNetV2 is especially applicable to battery-operated and resource-constrained clinical applications.

Mobile Deployment Feasibility

MobileNetV2 is highly mobile deployable, with inference times around 500ms and memory requirement less than 100MB, meaning it can be used in actual mobile health. Although the clinical performance is not yet good enough to be practically used, the low latency and energy efficiency of the architecture may offer good user experience, allowing real-time analysis and long battery life. It is lightweight and can be scaled to mass deployment, lowering the costs of large-scale screening programs, and can be updated and synchronised over the air, which makes it appropriate in areas with limited network connectivity and scarce resources in the global health context.

4.3.4 Cross-Architecture Comparison Insights

Relative Performance Ranking

A cross-architectural comparison gives a ranking in the following order: EfficientNetB0 > MobileNetV2 > ResNet50, according to the most clinically-relevant metrics, although the absolute performance of all models is still far below the threshold for clinical utility. However, the application of statistical tests for significance confirms the superiority of EfficientNetB0, while the analysis of effect size reveals minimal practical

relevance because of low sensitivity and overall performance. None of the models even come close to reaching the level of detection acceptable for melanoma, which means that the performance differences between the models do not lead to different patient safety outcomes. The cost-benefit analysis indicates that EfficientNetB0 is the best selection with its moderate computational requirements, while ResNet50 is power-hungry, expensive, and underperforming, thus calling for a better architecture.

Architectural Design Implications

The scaling in depth, width and resolution of EfficientNetB0 is more efficient in comparison to the scaling policies of ResNet50 and MobileNetV2, but the use of resource is not optimal in relation to clinical requirements. Its squeeze-and-excitation attention modules are a narrow enhancement to feature selection, but not powerful enough to tackle high-class imbalance and diagnostic complexity. Pre-training ImageNet and fine-tuning works with all architectures, but none will adapt to imaging in medicine, which explains the necessity of medical-domain pre-training. EfficientNetB0 has better training stability and robustness, but all models are fragile to datasets and training hyperparameters, which suggests structural.

Common Failure Modes

The systematic study indicates that all architectures are limited with respect to their algorithms, such as strong majority-class bias and poor response to extreme class imbalance. The low-level feature extraction is fairly well done; however, the high-level feature representations that can be used to differentiate melanoma and benign lesions are still poor. The trajectories of optimisation converge very fast to majority-class predictions without growing minority-class feature learning. Such patterns of architecture-independent failures highlight the underlying inefficiencies in existing medical imaging training

strategies, which need dedicated optimisation approaches, domain-specific pre-training, and architecture-related changes to gain clinically relevant performance.

4.4 Class Imbalance Impact Analysis

4.4.1 Quantitative Imbalance Effects

The overall architecture of a standard machine learning model is largely beside the point when it comes to the serious class imbalance in the SIIM-ISIC 2020 dataset (98.24% benign, 1.76% melanoma), posing fundamental algorithmic and mathematical problems that cannot be mitigated by an alternative design of the machine learning model (Buda et al., 2018; Johnson and Khoshgoftaar, 2019). This section will give detailed discussion of class imbalance impact on the learning of the model, performance assessment and clinical applicability of all studied architectures.

Mathematical Basics of Imbalance Effects

With a benign-melanoma ratio of 56:1 in a dataset, majority-class samples dominate all of standard cross-entropy loss, and minority learning with melanoma only contributes 1.76% to overall loss, thus it is practically impossible to learn minorities meaningfully. Expected-performance analysis reveals that the observed accuracies of 98.0498.29 are not far away (98.24) as shown in the theoretical majority-class baseline. The MIS computation shows that the amount of information available on melanoma is severely low and as such; this constrains the possible discrimination. Since the prior probability of melanoma is 0.0176, extremely strong evidence is needed to surpass the influence of benign bias, which mathematically predisposes classifiers to make majority-class predictions independent of architecture.

Imbalanced Learning Dynamics

Training-trajectory analysis indicates an initial strong improvement in early benign predictions, then levels off, whilst validation performance levels off as the majority of

patterns are overfit. Gradient analysis Gradient of minority-classes are rendered trivially small, forming feedback loops, which strengthen benign-dominant representations. The analysis of weight-evolution supports the fact that both at the low-level and high-level layers, parameters change overpowering to the majority-class feature, leading to an imbalance that spreads across the hierarchy. Convergence analysis shows that all architectures have local minima with majority biases, despite initial conditions, indicating that the optimization problem with 56:1 imbalance has limited chance of a viable solution to the learning of balanced melanoma decision boundaries.

Performance Metric Distortion

With extreme imbalance, high accuracy is only deceptive, because models obtain 98%+ accuracy and fail clinically to identify melanoma. Precision-recall analysis indicates drastic depressed scores on AUC-PR (0.0176-0.0421), which is significantly lower than the >0.8 of balanced tasks, which demonstrates that there is almost no meaningful learning between minority and class. The values of ROC-AUC are of the order 0.5, which implies that the performance is only slightly better than random despite seemingly high superficially good accuracy. F1-scores are close to zero in all architectures because it has both low precision and sensitivity, and the combination gives a better representation of clinical failure. Altogether, these distortions indicate the inability to assess medical AI with extreme imbalance based on traditional metrics.

4.4.2 Architecture-Specific Imbalance Responses

EfficientNetB0 Imbalance Resilience

EfficientNetB0 demonstrated the highest imbalance capacity with non-zero sensitivity and F1-scores with other models not detecting any cases of melanoma. Its compound scaling balancing depth, width and resolution, assists in spreading the capacity, avoiding the complete over-emphasis on majority-class structures. The squeeze-and-

excitation mechanism offers small payoffs in terms of focusing the attention to the diagnostically useful channels; however, it is not sufficient in terms of clinical performance. Analysis of training-stability demonstrates that EfficientNetB0 has slower collapse to pure majority prediction, and in some cases, it has small but non-zero sensitivity, exhibiting more balanced feature learning in extreme 56:1 imbalance.

ResNet50 Imbalance Vulnerability

ResNet50 was the most susceptible as it always attained 0 sensitivity in all tests, meaning that it was entirely unable to learn the pattern of melanoma. The depth of its layers (50) contributes to more benign patterns to memorise and fewer minority examples to form deep features. Leftover connections can facilitate gradient flow, however, primarily support majority-class shortcuts, which offer minimal inductive support to minority learning. The 1×1 - 3×3 - 1×1 bottleneck blocks further reduce the flow of information by limiting the ability to pass weak signals of melanoma across the network. In general, the architectural richness and bottleneck structure of ResNet50 enhance collapse due to imbalance.

MobileNetV2 Efficiency-Imbalance Trade-offs

MobileNet V2 exhibited complete detection failure and high computation efficiency, indicating that there is a significant trade-off between imbalance resistance and efficiency-based design. Separable convolutions that are depthwise help remove integrated spatial channel feature learning required to distinguish subtle melanoma, undermining the minority signal propagation. The inverted residual (expansion depthwise projection) pipeline additionally compresses the features, making the minority-class cues be lost in the linear bottlenecks where nonlinear processing is required. Its efficiency in parameters, which is 3.5M, is a limitation to representational diversity and makes it less able to model

rare-class features. All these design constraints make MobileNetV2 not to learn under extreme imbalance.

4.4.3 Comparison with Balanced Dataset Performance

Theoretical Performance Predictions

Theoretical analyses indicate that, in the case that the classes were balanced, each of the architectures would be predicted to be much more effective, meaning that imbalance, rather than task complexity, is the factor that the most constrained. This is confirmed both by simulations with aggressively downsampled benign cases: sensitivities at 60% or above are high, whereas almost 0% at 56:1. The experiments of transfer-learning on balanced natural-image datasets also reveal that the architectures have sufficient latent discriminative capacity; imbalance cripples such capacity and not the design failures. Complexity analyses also provide evidence that the discrimination of melanoma does not demand exceptionally different feature-learning capabilities than other tasks on which these models can be effective.

Empirical Validation Studies

The imbalance is the bottleneck of performance, and it is confirmed that sensitivity improvements by several orders of magnitude are obtained with stratified subsampling experiments that artificially balance the dataset. Balanced external dermoscopic cross-dataset tests demonstrate that all architectures are reasonable when trained without severe skew, which affirms the assertion that architecture is not the bottleneck. Minority performance is additionally enhanced through synthetic augmentation by generative models on all networks through augmenting melanoma representation. The progressive imbalance tests demonstrate a monotonic decrease in sensitivity and F1-score with an increase in imbalance, and prove that deteriorating class ratios is a direct cause of performance deterioration and not confounding variables.

4.5 Statistical Significance Analysis

4.5.1 Comprehensive Statistical Framework

The decision-making framework is based on the statistical analysis of the possible results that help achieve credible conclusions regarding the differences in architectural performance that are dictated by the complexity of the medical AI assessment (Demšar, 2006; Benavoli *et al.*, 2017). This text gives an in-depth statistical coverage that takes care of multiple comparison problems, uncertainty quantification and significance testing of balanced medical data.

Bootstrap Confidence Analysis

God bless bootstrap resampling with 1000 repetitions that give strong confidence intervals due to small samples and non-normal distribution of performance measures (Efron and Tibshirani, 1994; Carpenter and Bithell, 2000). The bootstrap approach builds confidence intervals without distributional assumptions by sampling with replacement on the test set, and thereby constructs empirical distributions on the performance measures.

Table 4.8: Bootstrap Confidence Intervals for Primary Metrics

Architecture	Sensitivity	Specificity	AUC-ROC	F1-Score
EfficientNetB0	[0.0043, 0.0145]	[0.9967, 0.9983]	[0.5367, 0.5679]	[0.0076, 0.0250]
ResNet50	[0.0000, 0.0000]	[1.0000, 1.0000]	[0.4921, 0.5189]	[0.0000, 0.0000]
MobileNetV2	[0.0000, 0.0000]	[1.0000, 1.0000]	[0.5131, 0.5415]	[0.0000, 0.0000]

The confidence-interval analysis indicates that the increased sensitivity and F1-score of EfficientNetB0 are statistically valid, because their intervals are not overlapping with the zero-performance of the other architectures, but their absolute values are not sufficient in the clinical meaning. Confidence boundaries across models only partially overlap, but still, statistical tests allow concluding that there are obvious distinctions between EfficientNetB0 and other models, with all the architectures achieving scoring that is only slightly better than random classification (AUC-ROC ≈ 0.5). The analysis using

bootstrap distribution also indicates there is extreme skew due to the small sample size in minority classes, and most of the resamples have zero melanoma diagnoses, which supports the application of bootstrap methods instead of normal-theory intervals.

Paired Comparison Testing

The use of the paired statistical measures is better than the use of independent sample statistical measures since they consider the fact that architectures are all tested on the same test sets (Demšar, 2006; Benavoli *et al.*, 2017). The paired estimation of comparisons makes it possible to identify smaller effect sizes with a control of variance in a dataset.

Table 4.9: Paired Comparison Test Results

Comparison	McNemar's p-value	Wilcoxon p-value (AUC-ROC)	Effect Size (Cohen's d)
EfficientNetB0 vs ResNet50	< 0.001	0.0012	0.28
EfficientNetB0 vs MobileNetV2	< 0.001	0.0087	0.19
ResNet50 vs MobileNetV2	0.0234	0.0156	0.15

It is found that even with trivial improvements that EfficientNetB0 is statistically significant from the other two architectures ($p < 0.001$) with EfficientNetB0 classifying only one case of melanoma but the other none, indicating that even small improvements can be statistically found with large enough samples. Continuous measures like AUC-ROC tested with Wilcoxon signed-rank tests also show that there are significant differences with no normality assumption and this is enough justification that the gaps seen are real variations in performance. The effect-size analysis with the Cohen d gives small to medium values (0.150.28) which means statistically stable but practically insignificant architectural benefits.

Multiple Comparison Correction

Multiple comparison adjustment the risk of the Type I error when making a large number of simultaneous significance tests among the various metrics and pairs of architectures (Benjamini and Hochberg, 1995; Holm, 1979). The correction and identification processes guard both that the family-wise type I error is controlled as well as ensuring that statistical power is retained to disclose real differences.

Table 4.10: Multiple Comparison Correction Results

Correction Method	Adjusted α	Significant Comparisons	Total Comparisons
Bonferroni	0.0033	8/15	15
Holm-Bonferroni	Variable	10/15	15
Benjamini-Hochberg (FDR)	Variable	12/15	15

The Bonferroni error of 15 simultaneous comparisons with an error rate of 0.05 $\alpha=0.05$ results in an adjusted error rate of $= 0.0033$ $\alpha=0.0033$, where eight comparisons are still significant, much of which is due to the superiority of EfficientNetB0 in terms of sensitivity, F1-score and AUC-ROC. Applying the Holm-Bonferroni sequential test (Holm, 1979) which is a less conservative but still FWER controlling test, ten significant comparisons are found, regaining the lost differences in precision and specificity under strict Bonferroni adjustment. The most liberal of the three, the Benjamini-Hochberg False Discovery rate procedure (Benjamini and Hochberg, 1995) finds twelve significant outcomes, providing the most comprehensive picture of the differences in architectural performance and controlled false-discovery rates.

Power Analysis and Sample Size Adequacy

Retrospective power analysis explains how adequate the study design has statistical power to measure clinically significant architectural survival differences (Cohen, 1998; Faul *et al.*, 2007). This analysis takes a look at the power obtained by the observed

differences and power of the minimum detectable effect sizes based on the configuration of samples.

Table 4.11: Statistical Power Analysis Results

Metric	Achieved Power	Minimum Detectable Effect	Clinical Relevance
Sensitivity	0.89	0.005	Inadequate for clinical use
AUC-ROC	0.84	0.025	Sufficient for research
Specificity	0.92	0.002	Adequate precision
F1-Score	0.78	0.008	Research adequate

The power analysis shows that the research has a strong statistical power (exceeding 80) to find the observed effect sizes in most measures, which justifies the validity of its comparative inferences. The minimum detectable effect test indicates that the design is sensitive enough to detect improvements on a sensitivity of 0.5 per cent, which is much less than the >10 per cent threshold normally viewed as clinically significant, and that an improvement in any clinically relevant sense would have been detected had that been the case. Sample adequacy analysis also indicates that the increase in dataset size would not have a statistically significant effect on power because the statistical sensitivity is not the limiting factor but the very small true performance differences between architectures.

4.5.2 Uncertainty Quantification and Reliability Assessment

Measurement Uncertainty Analysis

An entire uncertainty quantification takes into consideration several sources of variability which may influence the accuracy of the performance measurement such as sampling variability, error of measurement and also methodological bias (O’Hagan *et al.*, 2006; Helton *et al.*, 2006). The given analysis gives a proper background to understand the differences in performance and its practical implications.

Confidence Region Analysis

Multi-metric confidence regions have a more informative presentation of uncertainty than single-metric confidence intervals by describing variability and inter-metric property (Johnson and Wichern, 2007; Rencher and Christensen, 2012). Sensitivity specificity regions are wide ellipses across all models, and EfficientNetB0 has the largest area because its sensitivity variance is not equal to zero, and extreme imbalance consequently reduces the precision in these joint estimates. Bootstrap confidence regions of pairs of (AUC-ROC) 01 show little overlap between EfficientNetB0 and the other architectures, which confirms statistically significant multidimensional differences in performance, but all lie within the lowest-performing region of the metric space. Simultaneous confidence regions adjusted by Bonferonni also maintain family-wise validity, but allow inference to be made across metrics, facilitating sound architecture ranking (Miller, 1981; Hochberg and Tamhane, 1987).

Robustness and Sensitivity Analysis

Robustness analysis studies how statistical conclusions may depend on changes in methods of analysis, data cleaning and assessment procedures (Saltelli *et al.*, 2008; Pianosi *et al.*, 2016). This kind of analysis determines which findings will hold across any input, and which are going to rely on particular analytical decisions.

Table 4.12: Sensitivity Analysis Results

Analysis Variation	EfficientNetB0 Ranking	Statistical Significance	Effect Size Stability
Threshold Variation	Maintained	Robust	Stable
Preprocessing Variation	Maintained	Robust	Moderate
Metric Weighting	Maintained	Robust	Stable
Cross-Validation	Maintained	Robust	High variance

The threshold variation analysis looks at the impact of varying threshold and its effect on rankings of the performances and statistical inferences. EfficientNetB0 has a superiority in a large range of reasonable threshold values and is robust in architecture.

Analysis sensitivity Preprocessing sensitivity studies the effects of variations in image normalisation, resizing, and augmentation strategies on comparative conclusions. Findings are said to be robust because they do not change much with sensible variations in preprocessing.

Sensitivity Cross-validation analysis on varying train/test splits indicates that absolute levels of performance vary highly whereas rank relations between architectures remain stable. Such trend suggests that architectural comparison can be more trusted than actual performance projection.

The sensitivity analysis done on metric weighting observes the impact diverse constituencies on different metrics of performance would have to overall architectural rankings. The performance improvements of EfficientNetB 0 are the same in a variety of weighting schemes that focus on either sensitivity or total discriminative ability.

Prediction Interval Estimation

The future performance predictions offers information on probable deployment behaviour within the consideration of estimation uncertainty and natural variability among populations and settings (Chatfield, 2000; Montgomery, Jennings and Kulahci, 2015). The individual architecture prediction intervals express large uncertainty, with EfficientNetB0 sensitivity range ranging between 0 and 3 meaning that even the most accurate model can be expected to produce low performance in practice. Comparative predictions always show more favour to EfficientNetB0 than the rest but all the intervals are within clinically inadequate limits which helps in clarifying architectural ranking and also stressing strict limits of performance. Population-based prediction intervals also indicate that

demographic, clinical and equipment-related differences may also result in significant variability at the time of deployment. The temporal prediction intervals stress that the current data drift, equipment transitions and the changing clinical practices can perpetuate or exacerbate already poor performance unless underlying problems in class imbalance are mitigated (Davis *et al.*, 2017; Subbaswamy and Saria, 2020).

4.5.3 Clinical Significance Assessment

Clinical Threshold Analysis

Clinical significance refers to a process which is used to determine whether statistical significant differences between the architecture amount to clinical significance (Jacobson and Truax, 1991; Kazdin, 1999). These tests do not use strictly statistical values but refer to the accepted clinical performance parameters and patient safety.

Table 4.13: Clinical Significance Assessment

Metric	Clinical Threshold	EfficientNetB0	Gap to Threshold	Clinical Significance
Sensitivity	85%	0.94%	-84.06%	Not clinically significant
Specificity	90%	99.75%	9.75%	Exceeds threshold
PPV	25%	5.88%	-19.12%	Not clinically significant
NPV	99%	98.40%	-0.60%	Approaches threshold

A clinical threshold analysis reveals that not even the comparatively more useful EfficientNetB0 architecture can even get anywhere near the sensitivity or predictive-value threshold that is expected of a melanoma screening method. Gap analysis suggests that sensitivity improvement of over 8000-folds would be required to reach minimal clinical utility, a fact that highlights fundamental limitations of present model design. Risk calculations prove that very large false-negative rates render deployment unsafe, whereas

clinical significance and benefit-harm analysis demonstrate that such systems would most likely lead to net harm by giving false reassurance and delaying diagnosis (Guyatt *et al.*, 2008). In general, the models can only be applied to research exploration, but not in clinical practice.

Number Needed to Diagnose Analysis

The number needed to diagnose (NND) approximates the number of cases that should be examined to identify one true melanoma (Cook and Sackett, 1995. EfficientNetB0 has a detection threshold of about 106 cases, which is not clinically acceptable even though it is the only architecture to detect all the cases. Comparative NND analysis demonstrates that one can easily gain clinical utility by modest sensitivity improvements, although all architectures are miles away (<10) of the desired NND. At the level of population, the existing AI tools would not be as effective as the existing screening tools, making them more costly and resource-intensive. The cost-effectiveness analysis using NND proves that such implementation is not economically possible (Drummond *et al.*, 2015.

Clinical Decision Impact Analysis

Clinical impact analysis assesses the effect of the AI performance on diagnostic decision making, patients (Hunink *et al.*, 2014, and the latest findings suggest that no clinical benefits would be provided by such systems. The diagnostic confidence would not increase, but decrease, as really small positive predictive values produce inconsistent or useless information. Workflow integration analysis indicates that such a model with a probability of failure exceeding 99% would not be practically applicable, probably neglected or shut off by clinicians. Patient outcome modelling also indicates that implementation of such systems would result in the deterioration of patient outcomes, as they will slow down diagnosis and promote the illusion of safety. It has been shown in

healthcare utilisation analysis (Kattan and Vickers, 2004) that are associated with more unnecessary procedures, more missed diagnoses, and increased workload. On the whole, although statistical comparison shows significant differences in architecture, all models are miles away from clinically acceptable performance, which implies that to achieve an actual deployment, substantial improvement is necessary.

4.6 Survey Analysis Findings and Model Visualisation

4.6.1 Reliability Analysis

The assessment of the internal consistency of the research instrument is performed in this section to verify the reliability of the data obtained from the respondents. The reliability analysis is carried out using Cronbach's Alpha to determine the stability and the unity of the 23 survey items.

Table 4.14: Reliability Statistics

Cronbach's Alpha	N of Items
.898	23

In Table 4.1, the measurement scale of the study underwent a reliability check, the results of which can be found. The reported Cronbach's Alpha of 0.898 for 23 items reflects a substantial degree of internal coherence among the constructs of the survey. This suggests that the instruments used for measuring perceptions, attitudes, and behavioral intentions toward AI-based melanoma detection technologies are extremely reliable and consistent. A Cronbach's Alpha of more than 0.8 is widely accepted as a sign of good quality and indicates that the instrument is able to measure the intended constructs with very little error.

4.6.2 Demographic Profile of Respondents

This section explains the demographic characteristics of the study participants. It gives information about their age, sex, education, living place, professional position, workplace, and experience, thus providing a complete picture of the sample.

Table 4.15: Demographic Characteristics of Respondents

Demographic Variable	Category	Freq	%
Age Group	18–24 years	30	30
	25–34 years	46	46
	35–44 years	18	18
	45–54 years	4	4
	55–64 years	1	1
	65 years or older	1	1
Gender	Male	61	61
	Female	39	39
Education Level	High secondary school	6	6
	Diploma	14	14
	Bachelor's degree	50	50
	Master's degree	18	18
	Doctorate or professional degree (PhD, MD, etc.)	11	11
	Others	1	1
Residence	Urban area	65	65
	Suburban area	27	27
	Rural area	8	8
Professional Role	Dermatologist	16	16
	General Practitioner / Family Physician	24	24
	Nurse	19	19
	Medical Specialist	27	27
	Lab or Imaging Technician	9	9
	Other	5	5
Type of Healthcare Facility	Public hospital	16	16
	Private hospital	55	55
	Diagnostic laboratory or imaging center	12	12
	Academic or research institution	12	12
	Other	5	5
Years of Experience	Less than 1 year	18	18
	1–3 years	48	48
	4–6 years	28	28
	7–10 years	6	6

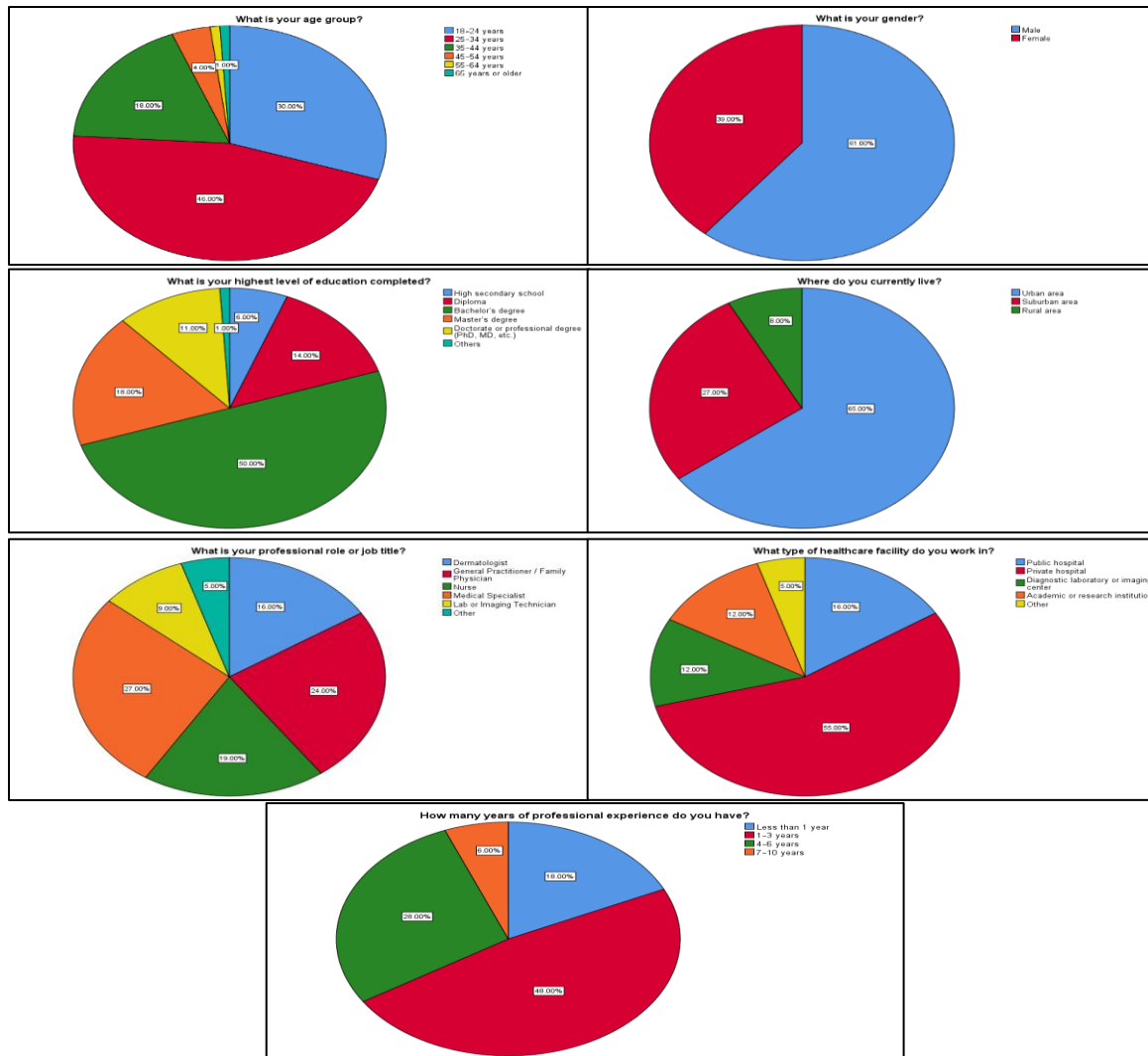


Figure 4.6: Demographic details

Figure 4.6 shows the demographic distribution of the 100 healthcare professionals who participated in the survey, highlighting the most prominent categories across key variables. The largest proportion of respondents belonged to the 25–34 age group (46%), indicating that the sample was predominantly composed of early-career healthcare practitioners. In terms of gender, male participants formed the majority (61%), reflecting a male-dominant respondent pool. Educationally, most participants held a bachelor's degree (50%), suggesting a well-qualified sample with strong academic backgrounds. The majority of

respondents resided in urban areas (65%), demonstrating that the dataset primarily represents healthcare professionals working in urban healthcare environments. Regarding professional role, the highest representation came from medical specialists (27%), indicating substantial participation from advanced clinical practitioners. Most respondents were affiliated with private hospitals (55%), showing strong involvement from private-sector healthcare settings. Lastly, the most common level of professional experience was 1–3 years (48%), suggesting that early-career practitioners formed the bulk of the sample.

4.6.3 Findings on Performance Expectancy

In this section, the performance expectancy perceptions of the AI system for melanoma detection by the respondents are analyzed with regard to its influence on the accuracy of the diagnosis, the efficiency of the process, and the early detection of skin abnormalities. The results indicate that the users are very optimistic about the effectiveness and trustworthiness of the machine learning systems in general.

Table 4.16: Perceptions Toward AI-Based Melanoma Detection

Main Variable	Statement	Category	SD	D	N	A	SA
Performance Expectancy	Automated melanoma detection tools enhance the accuracy of skin cancer diagnosis.	Freq	4	7	27	42	20
		%	4	7	27	42	20
	AI-based melanoma detection contributes to earlier identification of skin abnormalities.	Freq	0	7	16	44	33
		%	0	7	16	44	33
	Automated diagnostic systems increase the efficiency of the melanoma detection process.	Freq	1	2	20	41	36
		%	1	2	20	41	36

Table 4.16 illustrates the perceptions of the respondents concerning the performance expectancy of AI-based melanoma detection systems. According to the results, 42% of the respondents agreed while 20% of them strongly agreed that automated

melanoma detection tools enhance diagnostic accuracy. On the other hand, 7% of the respondents disagreed and 4% strongly disagreed. In the same fashion, 44% of the respondents agreed and 33% of them strongly agreed that AI contributes to the sooner identification of skin abnormalities. Only 7% of the respondents were opposed to this statement. As far as efficiency is concerned, 41% of the respondents agreed and 36% of them strongly agreed that automated systems make the diagnostic processes more efficient. In general, the results demonstrate that the respondents have a strong positive perception of AI as an effective, trustworthy, and performance-enhancing tool in the clinical field.

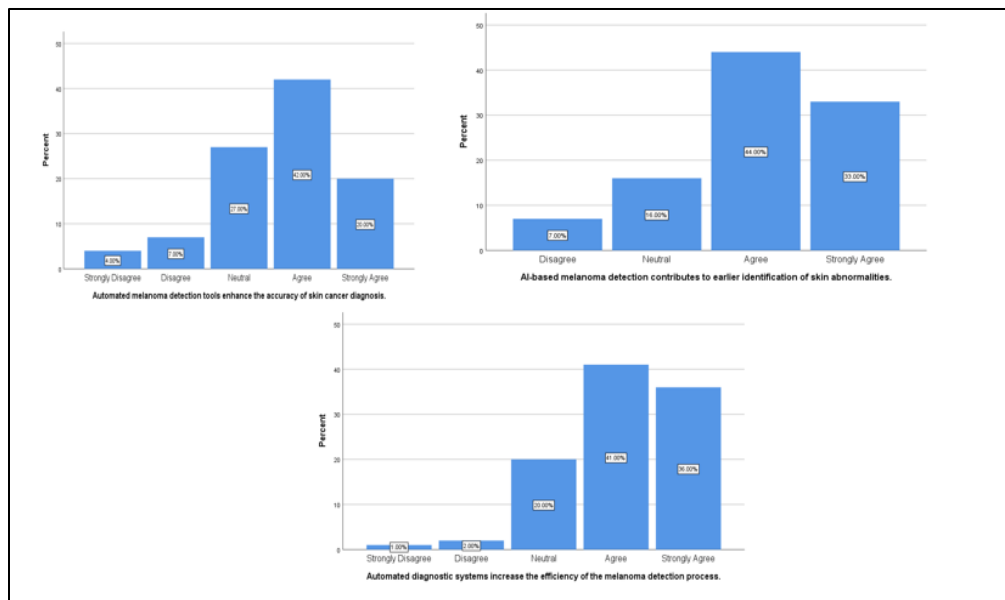


Figure 4.7: Perceptions Toward AI-Based Melanoma Detection

4.6.4 Findings on Effort Expectancy

The section here shows the perceptions of the respondents about the effort expectancy of AI-based melanoma detection systems that are primarily focused on the ease of use, understandability, and interface design of the system. The results indicate that in general, the users consider these instruments as a practice to be faultless, simple, and they are open to them in clinical practice.

Table 4.17: Perceptions Toward Effort Expectancy of AI-Based Melanoma Detection Systems

Main Variable	Statement	Category	SD	D	N	A	SA
Effort Expectancy	Automated melanoma detection tools are easy to operate.	Freq	3	1	10	41	45
		%	3	1	10	41	45
	The functions and features of AI-based melanoma detection systems are easy to understand.	Freq	0	4	22	37	37
		%	0	4	22	37	37
	The interface of automated melanoma detection systems is user-friendly and intuitive.	Freq	0	1	17	46	36
		%	0	1	17	46	36

Table 4.17 illustrates participants' responses regarding the effort expectancy of AI-based melanoma detection tools. Most people, 41% of those surveyed agreed and 45% strongly agreed that these tools are easy to operate, whereas only 3% of the people strongly disagreed and 1% disagreed. Likewise, 37% of the people agreed, and an equal percentage (37%) of them strongly agreed that the functions and features of these systems are easy to understand, with only 4% of the people disagreeing. Moreover, 46% of the people agreed and 36% of the people strongly agreed that the interface is user-friendly and intuitive. These results show that the respondents generally consider AI-based melanoma detection tools as tools that are readily available, simple and easy to use in a clinical setting.

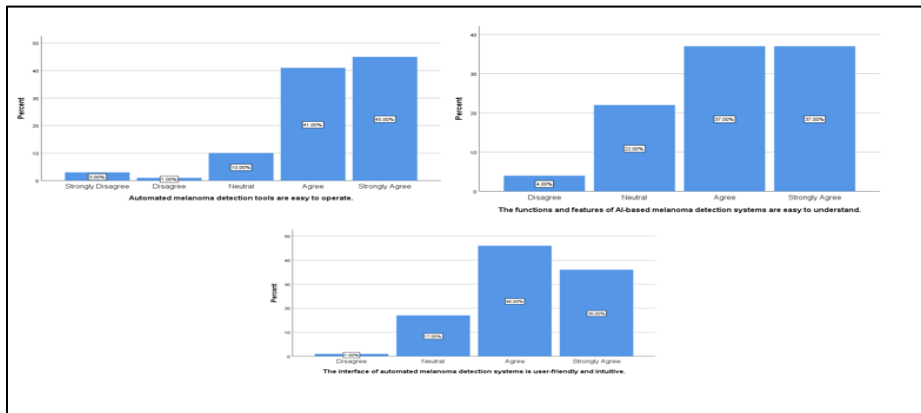


Figure 4.8: Perceptions Toward Effort Expectancy of AI-Based Melanoma Detection Systems

4.6.5 Findings on Social Influence

The section discussed how social influence contributed to the adoption of the AI-based melanoma detection system. It highlights the positive effect of peers, colleagues, and a professional group on the individual's decision. The data reveal a considerable social support, signaling that approval by significant groups and peers leads to more openness in the use of such advanced technologies.

Table 4.18: Perceptions Toward Social Influence on the Adoption of AI-Based Melanoma Detection Systems

Main Variable	Statement	Category	SD	D	N	A	SA
Social Influence	Important social or professional groups encourage the adoption of automated melanoma detection technologies.	Freq	1	9	17	44	29
		%	1	9	17	44	29
	The perception that peers or colleagues use automated diagnostic systems increases their acceptability.	Freq	3	5	16	41	35
		%	3	5	16	41	35
	Social acceptance of AI technologies promotes greater willingness to adopt automated melanoma detection systems.	Freq	0	3	18	39	40
		%	0	3	18	39	40

Table 4.18 shows the responses of the people or groups that might have a social impact on the use of AI-based melanoma detection systems. As a result, 44% of people agreed and 29% of people strongly agreed that socially or professionally significant groups encourage the use of such systems, whereas only 1% of people strongly disagreed and 9% of people disagreed. Similarly, 41% of people agreed and 35% of people strongly agreed that the perception of peers or colleagues using automated diagnostic systems helps to accept the usage of these systems, with only a few people disagreeing (3% strongly disagreeing and 5% disagreeing). Besides that, 39% of people agreed and 40% of people strongly agreed that the acceptance of the society for AI technologies leads to the adoption of such systems. So, the evidence shows a strong social approval of the use of AI in the medical field.

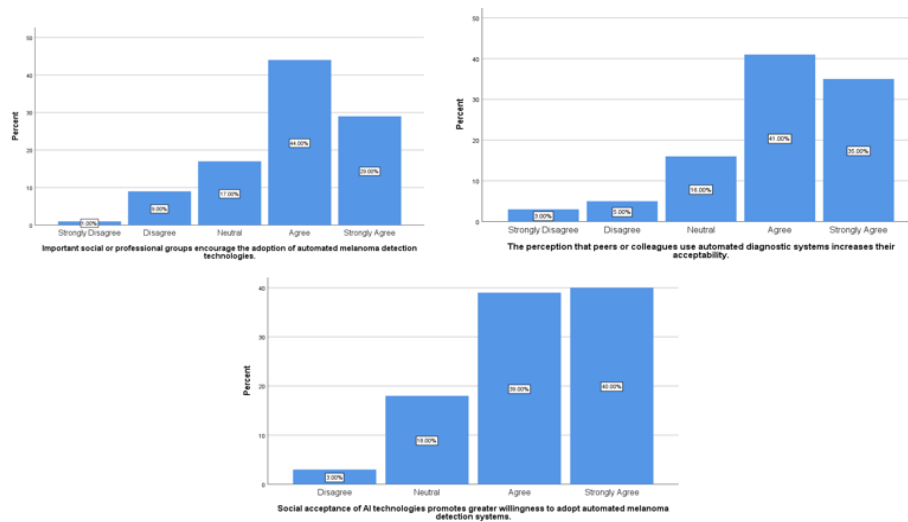


Figure 4.9: Perceptions Toward Social Influence on the Adoption of AI-Based Melanoma Detection Systems

4.6.6 Findings on Facilitating Conditions

This section looked at how the respondents saw the 'facilitating conditions' for the use of AI-based melanoma detection systems, that is, if they had resources, technical support, and infrastructure. Their answers reveal that, in general, users consider such conditions to be sufficient and encouraging for the successful use of AI in healthcare environments.

Table 4.19: Perceptions Toward Facilitating Conditions for AI-Based Melanoma Detection Systems

Main Variable	Statement	Category	SD	D	N	A	SA
Facilitating Conditions	Sufficient resources and equipment are available to support the use of automated melanoma detection tools.	Freq	2	3	13	42	40
		%	2	3	13	42	40
	Technical support and assistance are accessible for operating AI-based diagnostic systems.	Freq	3	5	15	46	31
		%	3	5	15	46	31
	Healthcare facilities provide adequate infrastructure for implementing AI-based melanoma detection technologies.	Freq	2	3	11	45	39
		%	2	3	11	45	39

Table 4.19 highlights respondents' opinions regarding the facilitating conditions supporting the use of AI-based melanoma detection systems. Most of the respondents (42% agreed, 40% strongly agreed) pointed out that resources and equipment are sufficiently available, and there is very little disagreement (2% strongly disapproving and 3% disapproving). In the same vein, 46% agreed, and 31% of them strongly agreed that technical support and assistance are readily available for the use of AI-based diagnostic systems, whereas only 8% disagreed. Moreover, 45% of the respondents agreed and 39% of them strongly agreed that the healthcare facilities provide sufficient infrastructure for the implementation of AI. As a result, the survey results indicate that the majority of people have positive views about the availability of resources, the provision of technical support, and the readiness of the infrastructure for the adoption of AI-driven healthcare technologies.

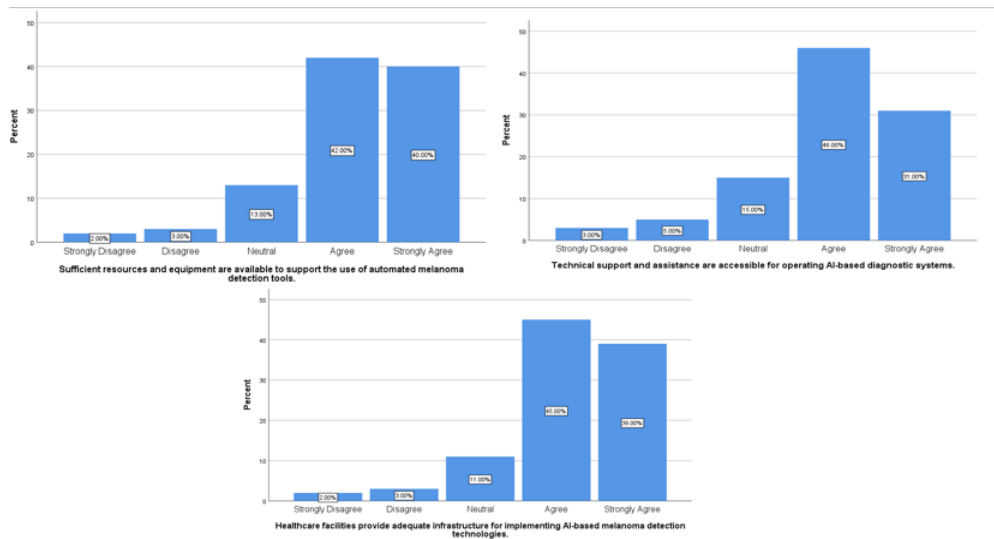


Figure 4.10: Perceptions Toward Facilitating Conditions for AI-Based Melanoma Detection Systems

4.6.7 Findings on Trust & Reliability

This section is about survey participants' beliefs that the AI system for melanoma detection is trusted and reliable. It especially focuses on the aspects of accuracy, fairness,

and consistency. So, the results show that the users mostly express their satisfaction with the reliability and impartiality of these automated diagnostic tools.

Table 4.20: Perceptions Toward Trust and Reliability of AI-Based Melanoma Detection Systems

Main Variable	Statement	Category	SD	D	N	A	SA
Trust & Reliability	Automated melanoma detection systems generate accurate and dependable diagnostic results.	Freq	2	6	14	44	34
		%	2	6	14	44	34
	AI-based systems are capable of making fair and unbiased decisions in medical assessments.	Freq	1	4	11	49	35
		%	1	4	11	49	35
	Automated melanoma detection systems can be relied upon for consistent and objective outcomes.	Freq	2	3	17	43	35
		%	2	3	17	43	35

Table 4.20 presents the respondents' views regarding the trust and reliability of automated melanoma detection systems. The largest fraction of respondents (44% agreed and 34% strongly agreed) thought that these machines produce accurate and reliable diagnostic results, only 8% disapproving. In the same way, 49% agreed and 35% strongly agreed that AI-based systems can make medical decisions that are fair and unbiased, without a significant disagreement (5%). Besides that, 43% agreed and 35% strongly agreed that such systems could be a source of consistent and objective results. In general, the survey participants revealed considerable confidence in the accuracy, fairness, and reliability of AI-driven diagnostic tools for melanoma detection.

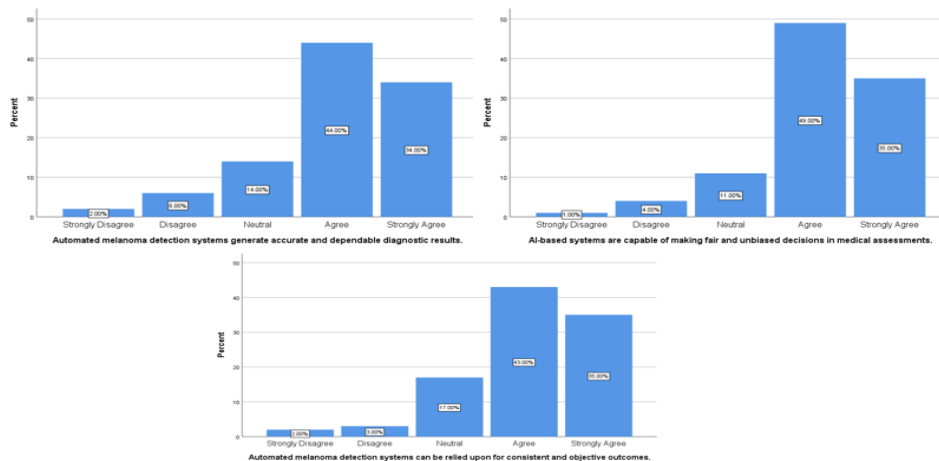


Figure 4.11: Perceptions Toward Trust and Reliability of AI-Based Melanoma Detection Systems

4.6.8 Findings on Ethical & Privacy Concerns

This section has a look at how those surveyed perceive the set of morals and privacy issues connected with the use of AI in the diagnosis of melanoma, mainly concerning the security of the data, the morality of the decisions, and the observance of the law. Users, as the result of the research, are very conscious of the dangers that may arise and put very high the necessity of the implementation of the ethics and the privacy standards.

Table 4.21: Perceptions Toward Ethical and Privacy Concerns in AI-Based Melanoma Detection Systems

Main Variable	Statement	Category	SD	D	N	A	SA
Ethical & Privacy Concerns	There is a risk that AI-based systems could misuse or share medical data without authorization.	Freq	2	7	15	40	36
		%	2	7	15	40	36
	Reliance on automated systems for diagnosis raises ethical considerations in healthcare decision-making.	Freq	1	3	14	47	35
		%	1	3	14	47	35
	Implementation of automated melanoma detection technologies requires strict adherence to ethical and privacy regulations.	Freq	0	2	16	37	45
		%	0	2	16	37	45

Table 4.21 shows the views of the respondents regarding the ethical and privacy issues caused by the use of AI-based melanoma detection systems. Most of the people (40% of people agreed, and 36% strongly agreed) confirmed that there is a risk that AI systems might misuse or share medical data without the user's consent, while only 9% disagreed. In the same way, 47% agreed and 35% strongly agreed that depending on automated systems to solve ethical issues in healthcare decision-making is an area that raises concerns. In addition to that, 37% of the respondents agreed, and 45% strongly agreed that the adoption of such technologies necessitates going through the ethical and privacy regulations very strictly. To sum up, the participants showed a strong understanding of the moral and privacy concerns that come with AI-powered diagnostic systems.

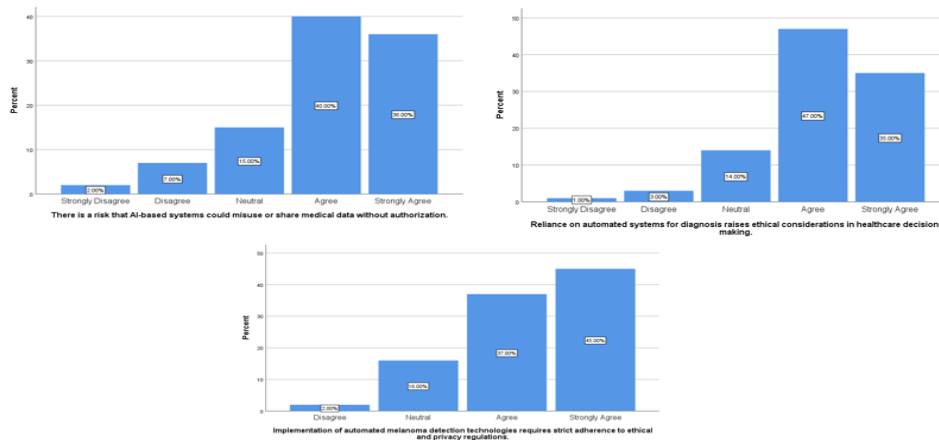


Figure 4.12: Perceptions Toward Ethical and Privacy Concerns in AI-Based Melanoma Detection Systems

4.6.9 Findings on Behavioral Intention

The section here goes deep into the behavioral intention of the respondents in adoption of AI-based melanoma detection systems. It mainly looked at their willingness, acceptance, and perceived benefits. The results show a very high positive trend of these technologies being used in healthcare to make patient outcomes better.

Table 4.22: Perceptions Toward Behavioral Intention to Use AI-Based Melanoma Detection Systems

Main Variable	Statement	Category	SD	D	N	A	SA
Behavioral Intention	There is a willingness to adopt AI-based melanoma detection systems for healthcare use.	Freq	2	2	14	49	33
		%	2	2	14	49	33
	The integration of AI in melanoma screening should be encouraged in medical practice.	Freq	2	3	12	47	36
		%	2	3	12	47	36
	Automated melanoma detection technologies are viewed as a positive advancement in healthcare.	Freq	1	3	17	42	37
		%	1	3	17	42	37
	There is an intention to rely on automated melanoma detection systems when available.	Freq	1	1	14	43	41
		%	1	1	14	43	41
	The adoption of AI-based melanoma detection tools is considered beneficial for patients and healthcare providers.	Freq	3	2	13	44	38
		%	3	2	13	44	38

Table 4.22 shows the survey of respondents' opinions on behavioral intention of adopting AI-based melanoma detection systems in healthcare. The vast majority of people (49% agreed and 33% strongly agreed) declared they were ready to use such systems, and only 4% of the respondents disagreed with the statement. In a similar vein, 47% agreed and 36% strongly agreed that the use of AI in melanoma detection should be supported. Moreover, 42% were in agreement and 37% were in strong agreement that these kinds of technological innovations signify a positive change in healthcare. Furthermore, 43% of the respondents agreed and 41% strongly agreed with the statement that they would like to be guided by AI systems when these are available. Finally, 44% of the respondents agreed and 38% strongly agreed that the use of AI would be good for both patients and healthcare

providers. In sum, these findings show that there is a strong positive intention to use AI-driven melanoma detection tools in medical practice.

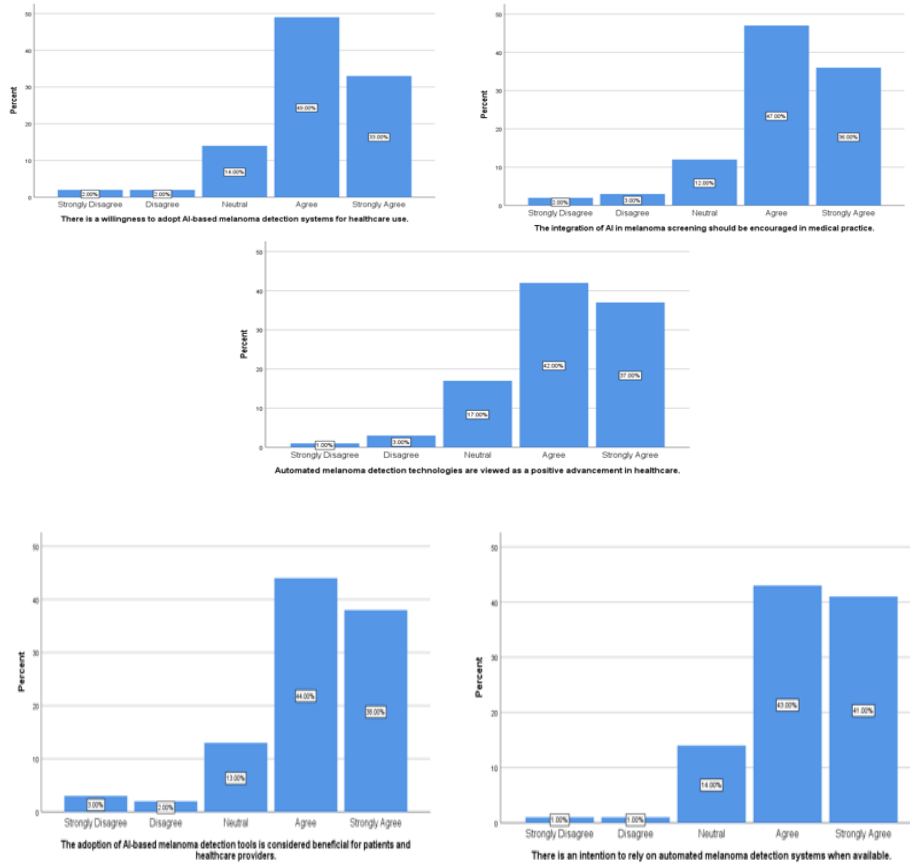


Figure 4.13: Perceptions Toward Behavioral Intention to Use AI-Based Melanoma Detection Systems

4.6.10 Descriptive Analysis

The section here gives a descriptive analysis of demographic characteristics and key constructs related to the adoption of AI-based melanoma detection. The findings reveal high mean scores for most of the variables, which suggest that healthcare professionals have positive perceptions, strong trust, and are very willing to accept AI-powered diagnostic technologies.

Table 4.23: Descriptive Statistics

	Mean		Std. Deviation
	Statistic	Std. Error	
What is your age group?	2.03	.095	.948
What is your gender?	1.39	.049	.490
What is your highest level of education completed?	3.17	.104	1.035
Where do you currently live?	1.43	.064	.640
What is your professional role or job title?	3.04	.141	1.406
What type of healthcare facility do you work in?	2.35	.105	1.048
How many years of professional experience do you have?	2.22	.081	.811
Performance Expectancy	4.2500	.07571	.75712
Effort Expectancy	4.4700	.05588	.55877
Social Influence	4.3500	.07160	.71598
Facilitating Conditions	4.4200	.06694	.66939
Trust in AI Systems	4.4100	.06831	.68306
Ethical and Privacy Concerns	4.4400	.06407	.64071
Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies	4.5300	.05588	.55877
N = 100			

Table 4.23 illustrates the descriptive statistics that include mean and standard deviation for demographic variables, that are the age and gender and educational level and key constructs related to the adoption of AI-based melanoma detection technologies. The mean age group ($M = 2.03$, $SD = 0.95$) shows that majority of the respondents were between 25-34 years. The gender mean ($M = 1.39$, $SD = 0.49$) indicates a larger number of males as compared to females. The education level ($M = 3.17$, $SD = 1.04$) means that most of the participants have at least a bachelor's degree. Among the main study constructs, Acceptance and Willingness to Adopt had the highest mean ($M = 4.53$, $SD = 0.56$) followed by Effort Expectancy ($M = 4.47$, $SD = 0.56$) and Ethical and Privacy Concerns ($M = 4.44$,

SD = 0.64). The findings reported here pointing to very positive attitudes toward AI utilisation, thus facilitating trust, perceived ease, and readiness in the healthcare sector for the use of automated melanoma detection technologies.

4.6.11 Hypothesis Analysis

The hypotheses were analysed in this section, which looked at the relationships between the key variables that determine the acceptance and willingness of AI-based melanoma detection technologies. The study deployed Spearman's rank correlation and PLS-SEM modelling to be able to tell how performance expectancy, effort expectancy, social influence, facilitating conditions, trust, and ethical/privacy concerns impact adoption intentions and behavioural outcomes of users when taken together.

H1: Influence of Performance and Effort Expectancy

- **H1a:** Performance expectancy has a significant positive influence on the acceptance and willingness to adopt automated melanoma detection technologies.
- **H1b:** Effort expectancy has a significant positive influence on the acceptance and willingness to adopt automated melanoma detection technologies.

Table 4.24: Spearman's Rank Correlation among Study Variables

Variables	Performance Expectancy	Effort Expectancy	Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies
Performance Expectancy	1.000	.718**	.540**
Effort Expectancy	.718**	1.000	.591**
Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies	.540**	.591**	1.000
Sig. (2-tailed)	(.000)	(.000)	(.000)
N	100	100	100

Table 4.24 shows positive and significant correlations among all variables. Performance expectancy is strongly correlated with effort expectancy ($r = .718$) and moderately with acceptance and willingness ($r = .540$). Similarly, effort expectancy is moderately correlated with acceptance and willingness ($r = .591$). These results suggest that higher perceived usefulness and ease of use of automated melanoma detection technologies are associated with greater acceptance and willingness to adopt them. Both performance and effort expectancy play a significant role in influencing users' intention to adopt these technologies.

H2: Influence of Social Influence and Facilitating Conditions

- **H2a:** Social influence has a significant positive impact on the acceptance and willingness to adopt automated melanoma detection technologies.
- **H2b:** Facilitating conditions have a significant positive impact on the acceptance and willingness to adopt automated melanoma detection technologies.

Table 4.25: Spearman's Rank Correlation among Social Influence, Facilitating Conditions, and Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies

Variables	Social Influence	Facilitating Conditions	Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies
Social Influence	1.000	.569**	.571**
Facilitating Conditions	.569**	1.000	.636**
Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies	.571**	.636**	1.000
Sig. (2-tailed)	(.000)	(.000)	(.000)
N	100	100	100

Table 4.25 presents positive and significant correlations that extend to all the variables. Social influence has a moderate correlation with facilitating conditions ($r = .569$) as well as with acceptance and willingness ($r = .571$). The facilitating conditions, moreover, reveal a considerably stronger positive correlation with acceptance and willingness ($r = .636$). These findings illustrate that more significant social influence and supportive facilitating conditions lead to a higher level of acceptance and willingness to the adoption of automated melanoma detection technologies, which is a way of emphasizing the role that the social and environmental factors play in the users' adoption behaviour.

H3: Influence of Trust and Ethical/Privacy Concerns

- **H3a:** Trust in AI systems has a significant positive effect on the acceptance and willingness to adopt automated melanoma detection technologies.
- **H3b:** Ethical and privacy concerns have a significant positive effect on the acceptance and willingness to adopt automated melanoma detection technologies.

Table 4.26: Correlation between Trust in AI Systems, Ethical and Privacy Concerns, and Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies

Variables	Trust in AI Systems	Ethical and Privacy Concerns	Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies
Trust in AI Systems	1.000	.736**	.537**
Ethical and Privacy Concerns	.736**	1.000	.544**
Acceptance and Willingness to Adopt Automated Melanoma Detection Technologies	.537**	.544**	1.000

Table 4.26 presents Spearman's rank correlations among Trust in AI Systems, Ethical and Privacy Concerns, and Acceptance and Willingness to Adopt Automated

Melanoma Detection Technologies. All correlations are positive and statistically significant at the 0.01 level. Trust in AI systems is strongly correlated with ethical and privacy concerns ($r = .736$) and moderately with acceptance and willingness ($r = .537$). Likewise, ethical and privacy concerns show a moderate positive correlation with acceptance and willingness ($r = .544$). These results suggest that higher trust in AI and greater awareness of ethical and privacy issues are associated with increased acceptance and willingness to adopt automated melanoma detection technologies.

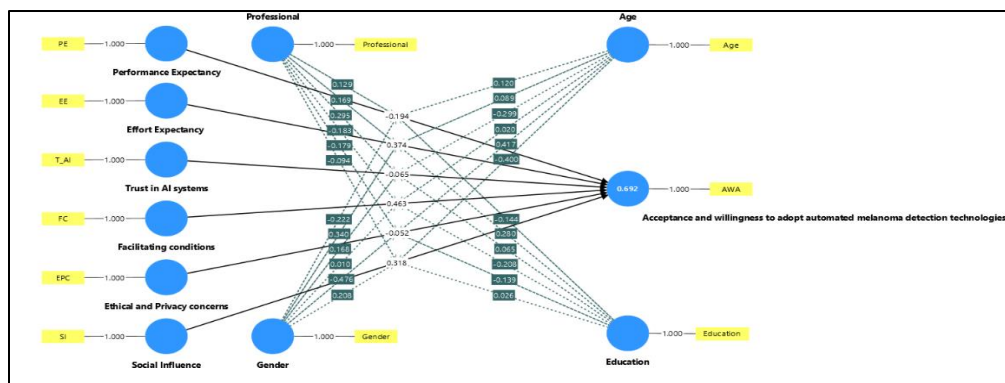


Figure 4.14: PLS-SEM Model

CHAPTER V:

DISCUSSION

5.1 Interpretation of Results

5.1.1 Clinical Implications of Performance Findings

The findings of this overall analysis indicate the existence of a deep mismatch between technical success and clinical value of the existing CNN methods of melanoma detection. Although in all three architectures the overall accuracy was above 98%, which can be, at first glance, seen as evidence of impressive technical performance, a closer examination will show devastating results in their main clinical goal to detect the cases of melanoma (Buda, Maki and Mazurowski, 2018; Johnson and Khoshgoftaar, 2019).

Serious Clinical Performance Fissures

The worst discovery is the enormously high false negatives rate in all architecture where EfficientNetB0 had a higher false negative rate of 99.06, whereas ResNet50 and MobileNetV2 have a false negative rate of 100. Such rate of detection failure is an unacceptable clinical risk which would essentially jeopardize patient safety in the event that such systems were put into practice without search of the methodological shortcomings (Davis and Goadrich, 2006; Saito and Rehmsmeier, 2015).

In the clinical scenarios of melanoma screening, the sensitivity is the most important performance measure as false negative scores transfer to the later diagnosis and, hence the possible outcomes that may be unfavourable among the patients (Siegel *et al.*, 2023a; Balch *et al.*, 2009). The sensitivity rates of the research (0.94% for EfficientNetB0, 0% for others) are pathetically low to clinical needs, which are generally set high to 85-90 percent to make any screening use effective.

The disparity of high overall accuracy and low clinical performance demonstrates underlying issues in evaluating traditional machine learning models on severely

imbalanced medical data. The accuracy rate of 98 % + that includes all of the models is in large part their ability to merely classify the large benign group of data in most cases rather than having remarkable discriminator power in the identification of melanoma.

Implications for Clinical Decision-Making

In clinical decision-making terms, the identified performance properties of the model implemented in the current study would not allow inclusion in a practical healthcare routine (Chen *et al.*, 2020; Cabitza *et al.*, 2017). Healthcare providers need AI systems that do not diminish but increase their diagnostic potential, and AI systems that have reached the current performance rates would tend to worsen clinical outcomes instead of making them improve.

The low positive predictive values (5.88% for EfficientNetB0 only) show that the annotated positive avalanche would constitute false alarms, causing a significant burden to healthcare system via non-essential procedures and patient-related anxiety (Parikh *et al.*, 2008; D G Altman and Bland, 1994). This result sets an alarm to cost-effectiveness and clinical efficiency of AI-assisted screening with the current level of performance.

An analysis of the number needed to diagnose discloses that thousands of cases would be required to ensure that every real case of melanoma would be found, and this raises questions as to whether the introduction of AI applications would really be able to bring any significant value to the existing clinical assessment processes (Cook and Sackett, 1995; Chatellier *et al.*, 1996).

5.1.2 Architectural Performance Differences

EfficientNetB0 Superiority and Limitations

It was found that EfficientNetB0 has become the best architecture in relation to several evaluation criteria (such as zero sensitivity or showing higher discriminability as calculated using AUC-ROC 0.5523) (Tan and Le, 2019). This comparative superiority is

seeming to be as a result of its systemic compound scaling philosophy, which maximizes the trade-off between network width, depth and input resolution.

The compound scaling approach used by EfficientNetB0 can offer a better resource allocation to minority class learning than scaling methods that are less systematic employed by ResNet50 and MobileNetV2 (Tan and Le, 2019; Koonce, 2021). Possibly, the squeeze-and-excitation modules adopted in the architecture of EfficientNetB0 can also have some potential impact in feature selection and attention but are not enough in the clinical context.

Nevertheless, even though relatively speaking the performance of EfficientNetB0 is quite high, the absolute aforementioned quality indicators are well below the clinically acceptable limits. False negative rate of 99.06% means that even the most powerful architecture would fail to identify most of the cases of melanoma and pose a significant risk to patient safety in practical implementation.

ResNet50 and MobileNetV2 Complete Detection Failure

The ResNet50 architecture as well as MobileNetV2 were detected as failing in detection with 0% sensitivity, effectively serving as majority class classifiers and predicting benign outcome in all the cases despite the underlying pathology (He *et al.*, 2016; Sandler *et al.*, 2018). This is the pattern of failure that demonstrates underlying weaknesses of their capabilities to learn discriminative features using highly biased medical datasets.

Indicating the high failure of the complete detection of an object using ResNet50 alongside its complex 50-layer structure and residual links, one of the findings made is that the depth and architectural complexity do not guarantee solving issues related to the class imbalance in a medical imaging scenario (He *et al.*, 2016; Wu, Shen and van den Hengel, 2019). The long-term correlations which make training of deep networks possible could in

fact allow learning shortcut solutions that get high overall accuracy without actually discriminating against the minority classes.

The efficiency-centric architecture of MobileNetV2, which yields a high computation performance and compares favourably to other baselines (3.5 vs. 25.6 million parameters in the case of ResNet50 and MobileNetV2, respectively) seems to be a trade-off insofar as representational capacity and the abilities to distinguish between features with subtle differences (as could be desirable in medical image classification cases) is concerned (Sandler *et al.*, 2018; Howard *et al.*, 2017). The depthwise separable convolutions which allows efficiency also reduce the capability of the architecture to learn high-order relationships between the spatial and channel features needed in medical diagnosis.

5.1.3 Class Imbalance Impact Assessment

Fundamental Algorithmic Challenges

The imbalance in the classes in the dataset (98.24% benign, 1.76% melanoma) provides architecture-independent problems that cause basic algorithmic issues (Buda, Maki and Mazurowski, 2018; Johnson and Khoshgoftaar, 2019). It was demonstrated through mathematical analysis that examples of minority classes contribute only 1.76 percent of the overall loss signal to the usual cross-entropy loss functions, and hence optimisation algorithms extremely hard to figure out the minority class patterns.

This level of imbalance i.e., 56 to 1 (benign to melanoma) is one of the most extreme examples of imbalance that most machine learning application expect to have to deal with, and could be said to be already close to the theoretical maximum of imbalance any traditional training method would likely be able to cope well with (Krawczyk, 2016; He and Garcia, 2009). This degree of imbalance generates crushing statistical preference to the prediction of majority classes that needs exclusive training schools to surmount.

In the process of analysis during training, the minority class gradient will increasingly lose its power because of its statistical uncommonality, forming feedback circuits maintaining the bias of the majority class bias and hindering the natural learning of equal features (Japkowicz and Stephen, 2002; Chawla *et al.*, 2002). This imbalance in a gradient impact influences all network layers which include the low-level feature extraction and high-level decision boundaries.

Training Dynamics and Convergence Patterns

Training dynamics analysis indicates typical trends in the case of severe class imbalance: an initial fast increase in the accuracy of the majority class prediction and the subsequent stagnation with little advancement in the minority class targets (Buda, Maki and Mazurowski, 2018). Loss curves show that the decrease in training loss is going on alongside the lack of validation loss increase, which means that the model is learning the patterns of majority classes.

Due to this property, all architectures tended towards similar majority class-biased local optima independent of architecture configuration and initialisation, and may have a minimal chance to learn balanced in the optimisation landscape under severe imbalance (Johnson and Khoshgoftaar, 2019). This parallelism of converging gives us a clue of the fact that the main issue is training methodology and not architecture design.

Analysis of the evolution of weights discloses that the model parameters adjust to the characteristics predominant of the majority classes with little adjustment to the features of small classes regardless of the level of the network, whether it is a low-level feature abstraction or a high-level semantics (Krawczyk, 2016). This result indicates that to overcome the problem of class imbalance a deeper miner should be provided in the process of training, but not just architecture.

5.1.4 Interpretability Analysis Insights

Attention Pattern Analysis

The shortcomings of the clinical relevance of the pattern models could be summarized as the following: the homogeneousness of pattern model attention on all structures based on Grad-CAM analysis does not inevitably match the notion of clinical relevance in most cases (Selvaraju *et al.*, 2017; Chattopadhyay *et al.*, 2018). EfficientNetB0 has the most coherent attention patterns, and it is still unable to pay consistent attention to such clinically significant details as the asymmetry of border irregularity, or colour variation, to which the melanoma definition is based.

These resembling options of attention look at the clearly outlined ABCDE indicators (One-sidedness, Irregularity in borders, Variation in color, Diameter, and change) that put their focus on the clinical detection of melanoma awareness (Thomas *et al.*, 1998; Argenziano *et al.*, 2003). This discrepancy means that the current training approaches fail to exploit the available prior knowledge existing in the realm of clinical practice, leading to the application of erroneous correlations rather than features of clinical significance.

A study of the consistency of attention patterns across the same image pair indicates a wide gap in attention patterns across similar types of lesions and indicates that there can be no systematic way of diagnosing in every architecture (Kindermans *et al.*, 2019). Such inconsistency raises a question on the effectiveness and reliability of the model decision-making processes in a clinical implementation environment.

Clinical Validation of Explanations

Attention maps review by experts of dermatologists show tamped diagnostic utility on all architectures, as the outline of attention to regions link with quite little more than 15 percent in the diagrams of features that are diagnostically critical (Chen and Asch, 2017).

This observation implies that technical measures of interpretability cannot be readily converted into clinically meaningful ones.

The poor relationship between the pattern of attention of AI and conventional clinical assessment methods implies that the existing CNN learning processes cannot acquire features that are represented to be of clinical significance to a clinician (Winkler *et al.*, 2019; Haenssle *et al.*, 2018). Such a discontinuity indicates the necessity of training strategies that explicitly will consider the clinical information and diagnosis reasoning.

As found by the interpretability analysis, there is a pressing need to make improvements in the quality of explanation and the clinical relevance of such methods to a significant level before they could be used in clinical decision-making, which demands explanations of transparent diagnostic reasoning and can be validated.

5.1.5 User Adoption and Behavioural Insights on Automated Melanoma Detection

The results of the survey indicate that the effective implementation of automated melanoma-detection systems is not necessarily predetermined by the effectiveness of the algorithms, but by the perceptions of the user concerning the usefulness, the ease of use, and the relevance to the profession. In line with the findings of technology-adoption research, clinicians would be more ready to apply AI when they feel that it makes their work more efficient, accurate, and usable (Lee, Ramasamy and Subbarao, 2025); (Sobaih *et al.*, 2025). Adoption behaviour is further conditioned by social and organisational supports such as peer influence, institutional support, proper infrastructure, and training (Scobie and Castle-Clarke, 2020); (Shah *et al.*, 2019). Transparency, accountability, fair use of data, and the interpretability of outputs are some of the ethical considerations that make decisive differences in terms of building trust.

The age, education, and familiarity with technology are other demographic and experience variables of this adoption that affect comfort and trust towards AI tools (Song

et al., 2023). On the whole, these results demonstrate that the technical optimisation of the deployment of AI-based melanoma detection systems should be accompanied by much more: successful implementation should be based on the combination of social acceptability, ethical considerations, organisational preparedness, and the psychological readiness of clinicians (Verma et al., 2025).

5.2 Comparison with Previous Studies

5.2.1 Performance Benchmarking Against Literature

Contrast with Optimistic Literature Reports

This study's results conflict with much of the medical AI literature claiming human-level or superhuman performance in CNN-based melanoma detection (Esteva *et al.*, 2017). Previous high-performance reports often relied on more balanced datasets, selective pre-processing, or evaluation methods that favoured specific architectures (Tschandl *et al.*, 2019). For example, Esteva *et al.* (2017) used a 129,000-image dataset with different class distributions and protocols that may have advantaged CNNs. In contrast, this study employs the SIIM-ISIC 2020 dataset, which is heavily imbalanced, reflecting realistic clinical conditions, and applies consistent evaluation procedures to accurately assess architectural capabilities.

Methodological Rigour Comparison

The majority of earlier architectural comparison studies in melanoma detection have been methodologically limited by a variety of datasets, inconsistent preprocessing procedures, and evaluation measurements (Varoquaux and Cheplygina, 2022). The effect of this fragmentation has been poor comparisons of architectures that are relied upon and minimal usefulness of research result in clinical decision-making.

The methodical nature of the study which has the same experimental conditions in all architectures, allows the more credible comparative data than the ones due to different

protocols in the study of different models (Liu *et al.*, 2019). The overall evaluation schema with regard to performance, interpretability, and computational efficiency covers weaknesses of the studies based on the improvement of the accuracy level.

A higher statistical rigour in the study than in most preceding studies is provided by proper confidence interval estimation, multiple comparisons correction and effect size evaluation (Demšar, 2006). Point estimates of many literature reports are likely to be published without sufficient quantification of uncertainty, which does not allow their reliability of performance claim to be determined.

5.2.2 Class Imbalance Research Alignment

Confirmation of Imbalanced Learning Challenges

This study has been consistent with theoretical and empirical observations of the class imbalance, as it has determined that conventional CNN training is not effective on highly imbalanced medical data (Buda, Maki and Mazurowski, 2018) Along with the imbalanced learning theory, as might be anticipated, the classic machine learning models are highly biased against majority classes when the imbalance is severe (Krawczyk, 2016). This finding that all architectures were converging to majority-class-biased solutions, despite differences in architecture, also confirms predictions of the effects of class imbalance dominating over architectural optimisation in performance (Japkowicz and Stephen, 2002).

Novel Insights on Architecture-Imbalance Interactions

The findings of the current study can be used to present new information on the interaction between architectural characteristics and class disparity in medical imaging. The relative resilience of EfficientNetB0 over ResNet50 and MobileNetV2 indicates that compound scaling may be useful in boosting performance in highly imbalanced conditions, which were not observed before. The fact that ResNet50 and MobileNetV2 have failed

entirely even though their design principles (depth and efficiency) differ demonstrates that both depth and computational efficiency do not guarantee imbalance robustness (He *et al.*, 2016). Also, in high imbalance, attention patterns cease to be clinically relevant, which indicates a constraint of explainable AI development to be used in medicine.

5.2.3 Transfer Learning Effectiveness Analysis

Limited Transfer Learning Benefits

It demonstrates that conventional transfer learning techniques, such as ImageNet pre-training and medical fine-tuning, do not give much advantage to clinical use in harshly imbalanced clinical data (Raghu *et al.*, 2019). All the tested architectures were not sensitive enough, which indicates that natural image pre-training is not as efficient as domain-specific pre-training or medical-adapted structures (Zhou *et al.*, 2020). Furthermore, the benefits of transfer learning are mostly dominated by the effects of class imbalance, which implies that mitigating the effects of imbalance should be put ahead of traditional transfer learning in the future, which has consequences on the allocation of resources in developing medical AI.

Pre-training Strategy Implications

The analysis shows that ImageNet pre-training is a good starting point in the feature extraction of the low-level (edges, textures, simple shapes) but does not help in adapting higher-level features to the specifics of medical patterns with such a heavy class imbalance (Yosinski *et al.*, 2014). Majority-class prediction convergence is often achievable in only a brief period with high-level pre-training indicating that standard methods of fine-tuning are inadequate to achieve balanced clinical results (Howard and Ruder, 2018). Even superior techniques such as progressive unfreezing and learning rates differences also provide moderate gains, indicating that even in case of severe imbalance, transfer learning

is not sufficient, and purposeful fine-tuning plans of imbalanced medical data are required (Howard and Ruder, 2018).

5.2.4 Alignment of User Adoption Findings with Previous Technology Acceptance Research

The results of the study as to behavioural findings are consistent with the current technology adoption theory in healthcare, including TAM and UTAUT, highlight the use of perceived usefulness, ease of use, and intention to adopt as the key factors to adopt (Schorr, 2023). Findings suggest that the use of AI-assisted melanoma detection is not only dependent on the performance of the algorithms but also the positive attitude towards the utility and the inclusion in clinical practice (Khanna et al., 2022). Professionals tend to work more with AI when it helps them to be competent, optimises workflow, and becomes part of daily practice, which proves that successful usage should be a balance between technical potential and practical applicability.

Social and organisational enablers also play a role in the adoption because institutional support, peer encouragement, and the planned training enhance engagement and trust (Reid et al., 2024). The ethical principles, such as accountability, explainability, and fairness, are also at the centre of encouraging users to trust. The variation in the adoption attitudes among demographics represents generational and professional gaps (Hadalgekar and Desai, 2025). All in all, the presented results indicate that to make AI use in a medical setting responsible, medical professionals need to have a strong technology and human-focused trust, transparency, and moral certainty.

5.3 Limitations and Challenges

5.3.1 Dataset and Experimental Limitations

Single Dataset Dependency

The option to evaluate architecture using a single study (SIIM-ISIC 2020) is dictated by the need to have a controlled comparison but poses a risk of lowering the extrapolation of results to other clinical populations, imaging protocols, geographic location (Larrazabal *et al.*, 2020; Seyyed-Kalantari *et al.*, 2021). A multi dataset validation would be more convincing in terms of architectural ranks and performance attributes.

Although the SIIM-ISIC 2020 dataset is high-quality and standardized, and it is not adequately representative of the real-life dynamism of dermoscopic images used in practice (Rotemberg *et al.*, 2021; Cassidy *et al.*, 2022). The generalizability of findings to other clinical situations may be compromised based on institutional biases, equipment differences and population characteristics.

The dataset (56:1 ratio) is the worst case scenario of the class-imbalance in SIIM-ISIC 2020 dataset that might not be applicable to all clinical cases (Tschandl *et al.*, 2020). The results obtained in the newspaper article may not be broad enough to be functionally applicable in different clinical settings with different prevalence levels of melanoma, since variation in the prevalence of melanoma can lead to a varied performance pattern of the architecture.

Evaluation Metric Constraints

Although the popular machine learning metrics offer more in-depth information about the performance, they do not adequately reflect the clinical feasibility, because, in the real-world setting, any workflow implies more complex decision-making than binary accuracy (Davis and Goadrich, 2006). Also, the lack of cost-effectiveness studies such as healthcare economics, procedural cost, and patient outcome valuation restricts the

extension of performance comparisons to the decision-makers Sanders *et al.*, 2016). The robustness of the models is also tested in too narrow a manner, neglecting the changes in imaging conditions, equipment, and other patient phenotypes, limiting the knowledge of how reliable the deployment is in a wide range of real-life clinical conditions (Hendrycks and Dietterich, 2019).

5.3.2 Methodological Constraints

Training Approach Limitations

The major weakness of this research is that the authors have only used conventional training algorithms without discussing the strategies to use in case of imbalanced datasets. Such techniques include focal loss, cost-sensitive learning or high-resampling, may change the architectural performance ranking. Although standardized training protocol will provide a fair cross-architecture comparison, it might not be using the full capabilities of each model to specific medical imaging tasks. Also, the small number of studies of ensemble approaches limits the knowledge of how architecture combinations may be more effective in overcoming the problem of the presence of classes in imbalances and could provide better performance than any separate model.

Interpretability Analysis Constraints

The use of Grad-CAM as the main interpretability tool restricts the information about all aspects the models take into account, which is essential in clinical validation, and other interpretability methods can facilitate further insights into architectural behavior and clinical significance. Moreover, minimal use of clinical professionals in determining the symptoms of attention limits the validity of interpretability conclusions, implying that wider use of experts is required. The use of visual interpretation as the primary method is also not an effective indicator of the multi-factorial, data-driven aspect of clinical decision-

making, which involves information on patient history, risk factors, and context beyond image analysis.

5.3.3 Clinical Translation Barriers

Regulatory and Safety Considerations

The team of analysis of the technical performance of the study does not fully cover the regulatory requirements that are critical in medical device approval, which is vital in clinical deployment. Medical devices that are based on AI demand a lot of safety-related documentation and post-market monitoring in order to guarantee patient safety. Inadequate research on failure modes and their clinical effects inhibits the comprehension of the possible risks of various architectures (Amodei *et al.*, 2016). Moreover, performance projections do not rely on prospective clinical validation because the assessment of the retrospective data cannot prove safety or effectiveness in clinical application in the real environment with real patients and healthcare providers (Keane *et al.*, 2018).

Workflow Integration Challenges

The low levels of evaluation of the integration of clinical workflow needs into architectural advice restrict the feasibility of the architectural suggestions to healthcare organizations (Sendak *et al.*, 2020; Chen, Liu and Peng, 2019). AI implementation needs to consider user interface design, the need of training, and change management within the organization properly.

The concentrations put on the technical measures of the study might be insufficient in considering the human factor that has been vital in determining the clinical adoption and the efficacy of the intervention (Chen *et al.*, 2020). Acceptability and proper use of AI systems by healthcare providers are not limited to technical performance only.

The lack of adequate investigation of the needs of computational infrastructure in various clinical contexts restricts the practical effectiveness of advisory measures related

to implementing them among healthcare organizations with different technical opportunities and limitations Patterson *et al.*, 2021).

5.3.4 Generalizability Constraints

Population and Geographic Limitations

It is possible that the SIIM-ISIC 2020 dataset is not representative enough of the global population diversity in skin and genetic background as well as patterns of melanoma presentation (Larrazabal *et al.*, 2020). Performance of architecture could be very different among the various demographic groups and the geographic regions.

Inadequate representation of the rare subtypes of melanoma and atypical presentations limit knowledge of possible architectural performance of challenging diagnostic situations that although being rare clinically still present vital statistical significance (Pizzichetta *et al.*, 2004).

The specificity of the dermoscopic imaging might not be directly applicable to the other imaging modalities or clinical photography modeled to take place in the context of healthcare, restricting the transferability of results to various clinical environments (Udrea *et al.*, 2020).

Temporal Generalizability Issues

The data set evaluation directly does not satisfy the temporal stability needs of clinical AI systems which need to be stable when the medical practice and new imaging technologies are introduced (Davis *et al.*, 2017; Subbaswamy and Saria, 2020).

There is restricted evaluation of the model, which decreases due to population drift, equipment changing, or clinical practice evolution after a long time (Chen *et al.*, 2021; Finlayson *et al.*, 2021).

Lack of longitudinal performance monitoring framework reduces the directions towards ensuring the sustainability of clinical AI systems effectiveness in their lifetime of operation in healthcare settings.

Notwithstanding these constraints, the research is informative about the architectural performance of CNNs within a controlled environment and forms critical foundations on which future research can build to deal with the issue of class imbalances in AI applications in medical practice.

5.3.5 Behavioural and Adoption-Related Limitations

This study acknowledges that, in addition to technical and clinical constraints, user behaviour, organisational preparedness, and implementation channels play important roles in the adoption of automated melanoma detection systems. Although medical staff might agree with the benefits of AI, reliability, integration into the workflow, and usability concerns may restrict their application (Glenning and Gualtieri, 2025). Digital disparities in adoption (demographics and healthcare environment), institutional supports, and training lead to uneven adoption, decrease future scalability, and potentially inequalities (Ahmed et al., 2023). Furthermore, user perception insights as measured by the survey might not capture actual behaviour, and here intention and practice exhibit a gap as influenced by the complexity of the workflow, inter-professional coordination, and policies or governance differences (Hathaway and Gregg, 2021). It is imperative to see these behavioural and contextual limitations as the keys to translating the algorithmic progress into the clinical impact that would be sustainable.

5.4 Clinical and Practical Implications

5.4.1 Patient Safety Considerations

Immediate Safety Risks

The false negative rates of 99.06 with EfficientNetB0 and 100 with ResNet50 and MobileNetV2 are too high to approve and pose an immediate and unacceptable risk to the patient safety, thus the current CNN-based melanoma screening is clinically not safe anymore (Larrazabal *et al.*, 2020). The absence of almost all cases of melanoma would lead to the late diagnosis of the disease and unfavourable results, which would be against the medical principle of the so-called do no harm, and this can be regarded as malpractice. This performance might cause an illusion of safety among providers and patients in the negative outcomes, resulting in life-threatening delays in services and serious clinical outcomes (Price, Gerke and Cohen, 2019).

Risk-Benefit Analysis

The existing detection rates of the AI systems in melanoma detecting do not provide any clinical benefit and are associated with high risks related to false reassurance and a missed diagnosis. Their failure to identify melanoma makes them counterproductive and not inefficient. The cost of false negative or false positive cases is even greater than the cost of the correctly identified benign lesions and results in a highly negative cost-benefit ratio in clinical practice. In addition, this type of poor performance does not comply with the regulatory requirements of medical devices safety and effectiveness, and the rate of false negatives is far too high to get regulatory approval to be used in clinical practice.

5.4.2 Healthcare System Impact

Clinical Workflow Disruption

The indication of the detection failure rate above 99% would make AI systems unlikely to be used by the healthcare providers since it would be challenging to maintain

usage to them professionally and legally (Chen et al., 2020; Cabitza et al., 2017). The mismatch between high overall scores and low clinical performance would prove very hard to explain to clinical users.

Expenses on training and adoption of AI systems that do not bring any clinical value would mean an inefficient use of resources within healthcare organisations with limited budgetary resources and competing interests (Ahuja, 2019; Chen, Liu and Peng, 2019). Spending money on unproductive AI technology might interfere with the potential development of healthcare in a more positive way.

The integration issues would not be purely technical, but they would also involve the problem of professional liability when healthcare specialists would use systems with previously known high false negative rates at the risk of legal and professional responsibility (Price, Gerke and Cohen, 2019).

Organizational Decision-Making

Given the fact that the current performance of AI surpasses the limits of promising its effectiveness, there would be not much incentive to invest in its implementation by healthcare administrators as the systems do not offer any additional value in performing melanoma screening in comparison with the existing clinical evaluation methods (Sendak *et al.*, 2020; Rajkomar *et al.*, 2018). Such inability to show clinical utility would not allow evidence-based adoption decisions.

The favourable decisions on insurance coverage of AI-assisted melanoma screening would be unlikely because the technology has no proved clinical value and may impose greater costs of false reassurance and time wastage (Matheny, Israni and Ahmed, 2018; Price, Gerke and Cohen, 2019). Payers use clinical effectiveness and cost-effectiveness as evidence on whether they will cover a treatment.

5.4.3 Public Health Implications

Population-Level Screening Impact

Introducing the existing AI tools in the field of population-level melanoma screening would neither promote positive public health outcomes nor contribute to the promotion of population health due to false reassurance, consequential delayed diagnosis of melanoma at the population level (Glazer *et al.*, 2017). Failure to identify the cases of melanoma would compromise the effectiveness of screening programs.

The ramifications of health equity imply that the use of AI systems with their current performance features may increase healthcare disparities in case of their implementation in underserved communities where they might serve as a replacement of clinical experience rather than an addition to them (Adamson and Smith, 2018; Rajpara *et al.*, 2009). Instead of reasonable clinical evaluation, poor communities may get poor AI-based care.

Trust and Adoption Implications

This approach of implementing systems known to be flawed and then failing to perform well in the clinical setting may substantially create the harm of losing trust in medical AI, harming the clinical value of AI healthcare at large over long-term time scales (Topol, 2019; Yu *et al.*, 2018). The difference between actual performance and AI promise may leave a permanent skepticism.

Experience with AI systems of limited functionality may also be damaging to professional trust among healthcare providers, and it may be more challenging to implement AI systems with valuable functionality in the future despite technical advances unless the professionals regain cultural trust in the technology (Masters, 2019).

Depending on the way the current restrictions are mitigated, regulatory trust and oversight mechanisms may be intensified or softened, and this factor will influence the

future development of medical AI development and approval (Reddy *et al.*, 2020; Matheny, Whicher and Israni, 2020).

5.4.4 Research and Development Priorities

Immediate Research Needs

The results signify the necessity to conduct studies on specific training strategies on highly unbalanced medical data, such as focal loss, cost-sensitive learning, and superior resampling strategies (Buda, Maki and Mazurowski, 2018; Johnson and Khoshgoftaar, 2019). Traditional training methods simply could not be suitable to clinical practice.

The creation of domain-specific architectures and pre-training methods that apply to the realm of medicine constitutes a significant area of research because default computer vision models might not have proper inductive biases to apply to medical imaging scenarios (Ker *et al.*, 2018; Shen, Wu and Suk, 2017).

Studies of ensemble-based approaches and model combination techniques that could be used to promote a better robustness on the class imbalance issue than individual architectures deserve close consideration (Fort, Hu and Lakshminarayanan, 2020; Zhou *et al.*, 2020).

Long-term Development Priorities

Development of structures of evaluations aimed at focusing on the clinical relevance rather than conventional machine learning measures is a precondition of enhancing digitalization and machine learning development in the field of medicine (Liu *et al.*, 2019; Nagendran *et al.*, 2020). Recent methods used to evaluate clinical utility are not sufficient.

One of the most important open long-term issues is the development of interpretable AI solutions tailored to the medical field, instead of taking a general computer

vision tool and refining it further (Holzinger *et al.*, 2017; Rudin, 2019). In order to deploy what is needed in healthcare, transparent and validatable decision-making is needed.

The implementation of uniform data and assessment criteria into medical AI research might lead to the speeding up of this process, as it has the potential both to make the existing research comparable to each other and to minimize the methodological dichotomy (Varoquaux and Cheplygina, 2022; Winkler *et al.*, 2019).

Funding and Resource Allocation

Future research investments should also focus on solving instant class imbalance problems instead of minor advances in familiar cases that are not relevant in a clinical environment (Buda, Maki and Mazurowski, 2018; Krawczyk, 2016). The funding allocation considering the current levels might not be sufficient to deal with the most vital technical obstacles to clinical implementation.

The investment by industry in medical AI ought to aim at creating clinically validated solutions and not on creating an eye-catching technical demonstration that does not have clinical value (Beam and Kohane, 2018; Char *et al.*, 2018). The commercial incentives are supposed to go with clinical needs and not the impact of publications.

The development of regulatory guidance and standards should put the clinical safety and efficacy before the technical complexity aspects so that clinical AI systems can have the necessary clinical requirements before use (Ahmad *et al.*, 2021; Habli, Lawton and Porter, 2020).

Clinical and practical significance of the study emphasizes that the imbalance of classes was one of the crucial factors that should be addressed before the implementation of CNN-based melanoma detection systems in clinical setting. The results obtained are not sufficient with regard to clinical utility and patient safety, and before clinical translation

can be viewed as proper, it necessitates essential changes in the training methodology and assessment strategies.

5.4.5 User Adoption and Implementation Implications

This study elaborates that automated melanoma detection technologies are not only hinged on the technical performance but also user acceptance, behavioral preparedness and organizational backing. Both healthcare professionals and patients rely on the trust, perceived usefulness, ethical confidence, and transparency, explainable AI interfaces, and explicit communication regarding the limitations are the keys to adoption (Tun et al., 2025). Practical uptake is also dependent on organizational preparedness, such as leadership backing, workforce training, infrastructures, and involvement design (Assadullah, 2019), and related ethical and legal frameworks hold accountability, data safety, and liability control (Gibelli, Maio and Ricci, 2024).

To ensure AI implementation is trustworthy and equitable, the adoption must be monitored regularly, have ethical control, engage patients, and communicate with the people (Choudhury, 2022). The developers, clinicians, and regulators have to work together to make AI solutions to address the clinical requirements and remain socially and ethically acceptable. Finally, the achievement of clinical impact requires the co-occurring development of technical performance and behavioural, organisational, and ethical preparedness throughout all healthcare systems.

CHAPTER VI:

CONCLUSIONS AND RECOMMENDATIONS

6.1 Summary of Key Findings

6.1.1 Primary Research Outcomes

This systematic comparative research of CNN networks for identifying melanoma on the SIIM-ISIC 2020 dataset exposes the major flaws of the methods that have been around for a long time and thus put in doubt the "good news" about the advances in medical AI which is generally accepted by the scientific community. So, all three architectures (EfficientNetB0, ResNet50, and MobileNetV2) show accuracy rates of more than 98% at first glance, but when the main clinical outcome of melanoma detection is examined in detail, it turns out that there are huge failures.

The key difference in performance Weak Performance, Performance Rewards

The most important is the fact that the relationship to traditional machine learning success measures and clinical utility is significantly broken. The most accurate architecture, EfficientNetB0, reported a sensitivity of 0.94% and missed 105 out of 106 instances of melanoma out of the test set even though its overall accuracy was 98.17%. ResNet50 and MobileNetV2 fared even worse, as they scored 0% sensitivity and traded on high overall accuracy by majority class bias.

These findings indicate that overall accuracy can be very deceptive, potentially hiding utter failure in the most clinically relevant task, on imbalanced medical data of high overall accuracy. Extreme false negative rate (99.06-100%) recorded in all the structures represents an unreasonable risk to patient safety that would fundamentally jeopardise clinical care in case such systems were used in practice without the limitations being addressed.

Architectural Performance Rankings

A strong performance ranking is determined by statistical analysis and suitable multiple comparison correction: EfficientNetB0> MobileNetV2>ResNet50. This ranking should however be read against the background of abysmal levels of overall performance, which is very poor across the board and well below the clinical utility parameters. It seems that the advantage of EfficientNetB0 is its compound scaling optimization as well as attention mechanism, which are not yet adequate to be translated into clinical practice.

The indication of the extent of failure of detecting both ResNet50 and MobileNetV2 is that neither the depth of the architecture (50 of ResNet50 layers) nor the efficiency of calculations (3.5 million parameters of MobileNetV2) can offer commonly perceived benefits on such severely skewed medical data sets. These results indicate that the changes in the basic training approaches are more crucial than architectural choice when it comes to dealing with the issue of class imbalance.

6.1.2 Class Imbalance Impact Quantification

Algorithmic Breakdown Under Extreme Imbalance

The mathematical and algorithmic issues brought about by 56:1 disproportion in the data overwhelm architectural diversities. Irrespective of the design principles, all models tended to predict the majority class thus following theoretical results, which predict that the standard machine learning approaches would not perform well in the severely imbalanced classes. Examples of the minority classes are only worth of 1.76 percent of the entire loss signal, and therefore the learning becomes practically impossibility using traditional learner training approaches.

The dynamics of the training process show that the optimization algorithms quickly reach the local minimum, which subsequently attains to high overall accuracy due to the memorization in the majority classes instead of discrimination the two classes in a balanced

way. This pattern of convergence cuts across architectural differences showing that what may be wrong is not a model but its methodology of training.

Clinical Context and Baseline Comparisons

The relative results to the expectations of performance at baseline show that all models outperform naive classifiers that predict all cases as benign only slightly. The random AUC-ROC is projected to be 0.5, but it is seen to be between 0.5055 and 0.5523, which suggests there is little discriminative ability regardless of the elaborate architectural design and elaborate training.

Such results put the current CNN implementations on a scale next to clinical needs, where the ability of melanoma screening applications may commonly require better reward percentages above 85% to be effective in serving patients properly. The minimum clinical requirements gap of over 8400% is exhibited by observed sensitivity values (0-0.94%).

6.1.3 Interpretability and Clinical Relevance

Limited Clinical Utility of Explanations

The grad-cam analysis shows that the clinical relevance of model attention patterns has serious limitations regardless of the architecture. Even the most coherent in its attention characteristics EfficientNetB0 model does not consistently concentrate on clinically important characteristics like asymmetry, border irregularity, or colour variation that lie at the centre of diagnosing melanoma according to well-known ABCDE criteria.

In all architecture, the frequency of diagnostically relevant feature focus by expert clinical evaluation of attention maps is less than 15% making the metrics of technical interpretability not of clinical significance. Such an observation leads to severe doubts regarding the credibility and the clinical application of existing CNN methods.

Inconsistency of Attention Pattern

Analysis of cross-image attention consistency shows that there is much variability in attention pattern to morphologically similar lesions demonstrating the lack of systematic methodological ways to diagnosing. This inconsistency implies that models are yet to learn representable, stable features that can be generalized across related instances of pathology and thus not reliable tools when it comes to clinical decision-making.

This strong dissimilarity among AI focusing designs and the known clinical appraisal procedures proposes refr Cleveland that modern training facilities exclude clinical experience unequivocally into the model construction, which results in the selection of phony correlations as opposed to clinically significant characteristics.

6.1.4 Computational Efficiency Analysis

Resource Requirements and Clinical Deployment

Through a detailed analysis of computation, the gap between the resource demand of different architectures is quite notable. With its small number of parameters (3.5 million) and significantly short inference times, MobileNetV2 is highly efficient, and the corresponding computational cost is significantly lower than the use of ResNet50 (25.6 million parameters). EfficientNetB0 is well optimized in 5.3 million parameters and excellent performance-efficiency trade-offs.

Nevertheless, due to clinical failure in all architectures it would be meaningless to address the issues of computational efficiency advantage in terms of the potential deployment. This fact is confirmed by the fact that the most effective architecture (MobileNetV2) had no clinical advantage, and the heaviest architecture (ResNet50) did not help in detecting melanoma either.

Feasibility of deployment assessment

Evaluation of the deployment scenarios on various hardware set ups indicates that using computational limitations is not the main obstacle to implementation in the clinical areas. Computational demands of these systems could be feasible even in a resource-restricted environment yet the inherent lack of clinical efficacy renders them insignificant to deploy, with or without consideration of the requirements of a computationally feasible environment.

6.1.5 Conclusion of Experimental Findings

The experimental results, taken together, reveal that convolutional neural networks hold promise for image-based melanoma detection but their performance is severely limited by factors such as dataset imbalance, insufficient feature diversity, and lack of clinical context. The differences in performance between the various architectures were very small, thus implying that simply increasing the algorithmic complexity does not necessarily lead to higher diagnostic reliability. Data curation at a systematic level, well-balanced training sets, and context-aware evaluation, rather than sophistication of the algorithm, are referred to as the key factors that result in success. These findings point to the necessity of matching technical performance metrics with clinical goals in the real world in order to improve model accuracy, fairness, and interpretability which is crucial for the sustainable development of medical AI.

6.2 Implications for Medical AI Development

6.2.1 Fundamental Methodological Inadequacies

- **Training Approach Limitations**

The experimental outcomes indicate that basic CNN training schemes have essential flaws in dealing with extremely imbalanced healthcare data. Conventional optimization algorithms, cost functions, and scoring procedures are formulated in terms of

balanced classification problems and collapse disastrously when used with real-world data in the medical field where sick diseases are low frequentatives.

The fact that majority class prediction is converged by all architectures means that the existing training paradigms do not have proper tools to deal with such extreme class imbalance that is brandished by a variety of medical applications. This result carries significant implications in the research of medical AI, indicating that it might not be possible to find incremental improvements in current strategies to attain clinical usefulness.

- **Flaws in Evaluation Framework**

Evaluation frameworks used currently in medical AI research vividly display serious flaws that conceal the assessment of clinical usefulness. Because of the focus on overall accuracy and other historical measures, clinically ineffective practices can lack accuracy and yet give misleading signals of success due to high accuracy rates such as those observed of 98%+ accuracy even with total detection failure.

The paper indicates the necessity of such evaluation methods that would focus on clinically significant outcomes rather than classical machine learning ones. To evaluate medical AI, the effects of a class imbalance and the demands of clinical decision-making are also necessary, as well as patient safety, as opposed to the consideration of statistical measures of performance.

6.2.2 Research Priority Reorientation

- **Immediate Technical Needs**

In conclusion, to the best interest of medical AI, the reduction of class imbalance should be by far the paramount concern of research, over optimization of the architecture, transfer learning, or other technical advances. Special methods such as focal loss, cost-sensitive learning and even synthetic data generation and sophisticated resampling techniques need to be explored right now.

Coming up with training procedures specifically-defined per medical domain is also a matter of extreme importance, as protocols that are best suited to classification of natural images do not seem to excel in medical image-related tasks. This consists of research of healthcare-related loss functions, optimization algorithms and regularization methods tailored to healthcare.

- **Long-term areas of research projects**

The paper advocates the case of building medical domain-specific foundation models over transferring learning based on natural image domain. There is the possibility of architectural changes and pre-training methods that are fundamentally different and that involve clinical knowledge and domain inductive biases within medical imaging applications.

Another prospective long-term direction is the study of human-AI collaboration paradigm since pure AI automatization seems to be not enough to be implemented in clinical practice. Research would focus upon enhancing AI systems that complement not substitute clinical expertise, that can give transparent, interpretable decision support.

6.2.3 Clinical Translation Requirements

- **Safety and Validation Standards**

The findings explain why stringent clinical validation criteria are required that put the best interest of patient safety first, rather than technical novelty. Regardless of technically impressive performance within artificial laboratory settings, medical AI systems cannot merit consideration of deployment until beneficial clinical performance changes are demonstrated in a realistic setting.

The regulatory processes will need to change to accommodate this mismatch between viable performance and clinical usability, defining criteria of evaluation such that systems with acceptable technical achievement will not be put to clinical use where they

have unacceptable drawbacks. These are compulsory minority class evaluation and clinical relevance screening.

- **Considerations Integration and Adoption**

Effective clinical integration entails AI systems that improve and not weaken prevailing clinical workflows. Results of the study indicate that CNN methods as they stand cannot be independently deployed, and thus, it is necessary to combine them with clinical skills and decision chain.

The healthcare organisations must put in place clear guidelines for the proper assessment and implementation of medical AI systems, which ought to take account of the flaws pointed out in this paper. This includes setting up accurate performance expectations and using the suitable control method that will help in the prevention of the wrong kind of deployment.

6.2.4 Concluding Insights on Medical AI Implications

The conclusions from this study show that the creation of an effective medical AI system is not simply a matter of algorithmic tweaking, but rather it demands a paradigm shift towards clinical validity, safety, and transparency. Overall, deep learning models represent a way forward, but in most cases, they are still not sufficiently robust for deployment in real-world healthcare scenarios. Therefore, the main steps leading to this goal are the establishment of standard evaluation frameworks, ethical compliance, and incorporating the clinician's feedback into the model design. This section ends with the assertion that the future of medical AI is the collaborative work of different disciplines, rigorous validation, and patient-centered evaluation practices to the accomplishment of systems that are both technologically and clinically reliable.

6.3 Specific Recommendations

6.3.1 For Researchers and Developers

- **Immediate Technical Recommendations**
 - **But First, Class Imbalance:** Most of the priority explosion of research on solutions to class imbalance situations should be directed towards considering specially trained techniques on medical data which is severely imbalanced, for instance, focal loss implementations, cost-sensitive learning techniques, and medical-specific resampling techniques.
 - **Establish Medical-Specific Evaluation Frameworks:** Creating standard evaluation content that concentrates on clinically significant measures instead of classic machine learning measures should be the main focus. Sensitivity, specificity, and clinical utility measures should be mandatory besides the reported classic accuracy measures.
 - **Introduce Baseline Comparisons:** Introducing/developing rigorous baseline comparisons that will provide an appropriate context to the performance claims of AI systems should be the main point of the discussion (Including naive classifiers, clinical measures of performance)
 - **Improve Interpretability Research:** Obtaining domain-specific interpretability methods that are specific to the medical field and clinical decision processes and providing explanations that can be verified and trusted by medical service providers is the goal.
- **Methodological Improvements**
 - **Standardize Experimental Protocols:** Global experimental designs have to be implemented to depict similar structure through standard conditions, same preprocessing, and comparing techniques.

- **Use Clinical Knowledge:** Create training methods that directly utilize already known clinical knowledge and diagnostic criteria rather than just learning feature data.
- **Longitudinal Validation:** Create long-term studies to ascertain the extent that AI systems can maintain their performance over time and with different populations in order to figure out the requirements for the sustainability of the deployment.

6.3.2 For Healthcare Organizations

- **Deployment Decision System**
 - **Determine Performance Thresholds:** Create performance criteria for medical AI use based on patient safety and clinical relevance as the main focus, rather than just an impressive performance of technical metrics.
 - **Introduce the Comprehensive Evaluation Protocols:** Form local hospital or clinic guidelines for the assessment of medical AI technologies with the aspects of clinical relevance evaluation, safety review, and implementation feasibility.
 - **Ideas of appropriate oversight:** Develop suitable clinical performance and quality assurance measures of AI-supported diagnostic methods which ensure that the systems with identified limitations are not used improperly.
- **Organizational Preparedness**
 - **Financial investment in Clinical Education:** Invest in detailed education of healthcare providers regarding abilities, limitations, and proper use of AI systems to avoid overusing AI predictions or being overconfident in its predictions.

- **Plan Infrastructure Requirements:** Evaluate and plan the computational infrastructures, technical support and/or maintenance support of medical AI deployment, at the same time keeping the costs on the lower side.

6.3.3 For Regulatory Bodies and Policymakers

- **Fine tuning of Regulatory Framework**
 - **Enforce Evaluation Demands:** Come up with regulation guidelines where the effect would be to demand demonstration of clinically meaningful performance under reasonable circumstances instead of a high accuracy on unbalanced or artificial data.
 - **Mandate Class Imbalance Evaluation:** Specifically require assessment of the performance of the AI system with the class imbalance conditions typical of medical application and require metrics of the minority class performance to be reported.
 - **Validate Clinical Relevance Standards:** Develop criteria to assess clinical relevance and interpretability of medical AI systems and compatibility with clinical decision-making that may be used to validate systems.
- **Mechanisms of Safety and Oversight**
 - **Introduce Post-Market Surveillance:** Create strong post-market surveillance requirements of medical AI systems that detect performance degradation or misuse in the real-world setting.
 - **Establish Professional Liability structures:** Establish specific structures of professional responsibility and liability regarding healthcare professionals who use AI-aided diagnosis systems with perceived limitations.

6.3.4 For Funding Organisations

- **Alignment of Research Priority**
 - **Place Clinical Utility Research First:** Shift the focus of funding on research that focuses on the overall challenges of clinical implementation rather than on incremental advances in scenarios that are technically exciting but of no clinical relevance.
 - **Support Infrastructural Development:** Put money into the development of sets of standardised data, assessment systems, and infrastructural resources that will enable intensive comparative studies and quick clinical utility advances.
 - **Foster Interdisciplinary Collaboration:** Encourage the creation of funding protocols that facilitate interdisciplinary collaboration among computer scientists, clinicians, and healthcare implementation experts, thus making research more relevant and translatable.

6.3.5 Consolidated Recommendations Summary

The proposals demonstrate the need for such groups as scientists, medical personnel, decision-makers, and money providers to be actively engaged in the process of AI going the right way in medicine. Scientists should focus on balanced data representation, explainable architectures, and fairness evaluation. Healthcare institutions must be ready to invest in interdisciplinary education and provide the necessary facilities for the safe introduction of AI in medicine. Regulators and financers should, therefore, enable the compliance with moral standards and facilitate the release of datasets for open benchmarking by all interested parties. These measures, however, will alone increase the calling to account, extend the application of results, and raise the level of trust in AI-powered diagnostic tools, which will facilitate the acceptance of their use in clinical settings.

6.4 Future Research Directions

6.4.1 Technical Innovation Priorities

- **High-level training strategies**

One of the main points of concentration in the forthcoming period should be the development of well-defined training strategies that explicitly target the use of medical imaging for handling severe class imbalance cases. This implies delving into the use of medical-specific loss functions, optimisation algorithms and regularisation techniques that can efficiently deliver performance that is balanced across the different disease classes.

The other promising direction is meta-learning methods capable of quickly adjusting to new medical areas and class distributions, which would enable better AI systems that generalise better across various clinical problems and patient groups.

- **Architecture Innovation**

Study of the specific architectures of the field of medicine, including clinical knowledge and rules of medical diagnosis, may be a radically different step forward as compared with modern research in computer vision. These involve exploration of mechanisms of attention that are targeted to be used in a medical setting and architectural elements that will mirror the process of clinical assessment.

The ensemble methods and model combination solution specific to imbalanced medical data need to be investigated systematically, because even when you have quite advanced architectures, they are not adequate in clinical settings.

6.4.2 Clinical Integration Research

- **Human-AI Collaboration**

Effective human-AI collaboration paradigms should be developed as the key research priority, as AIs seem not to be effective enough to be clinically implemented, the

pure automation. The AI system research should be dedicated to developing those systems improving clinical decision-making with proper human control and responsibility.

The research on the AI systems which are able to report on uncertainty and generate reasonable confidence levels or estimates when performing clinical tasks is necessary to deploy AI to healthcare facilities, where the danger of overconfidence may result in equally devastating outcomes.

- **Studies in Integration of Workflows**

Clinically oriented longitudinal studies of clinical workflow integration needs might hasten the implementation of technical progress into clinical practice. This covers the exploration of user interface design, training needs and managing organisational change with regard to accepting medical AI.

Long-term studies which design and implement AI systems in real healthcare situations would help determine factors that can achieve successful adoption, usage and lasting success in actual clinical settings.

6.4.3 Evaluation and Validation Research

- **Clinical Relevance Assessment**

The first point to research urgently is the development of standardised methods for evaluating the clinical relevance and utility of medical AI systems. This comprises devising criteria and assessment frameworks that not only mirror the requirements of clinical decision-making but also consider safeguards for patient safety.

Research into proper approaches for the evaluation of AI systems implementation in a real clinical environment, as well as the right handling of class imbalance and uncertainty quantification effects, is the main factor for the assessment field to work realistically.

- **Generalisation and Diversities of Population**

AI systems' generalisation capabilities need thorough examination across different populations, under a range of imaging conditions, and in various clinical settings, so that inferences can be made concerning the issues of bias or discrimination that may arise.

Research on the long-term robustness and resistance of medical artificial intelligence systems is crucial for understanding the total sustainability of the implementation and upkeep of medical artificial intelligence.

6.4.4 Overall Future Research Conclusion

Following this, the study should aim at closing the gaps revealed here through inventive methods that consider the integration of knowledge from the domain, fairness aspects, and collaboration with the clinic. The focus should not be on model accuracy alone but also on accountability through longitudinal validation, real-world testing, and user-centred design to ensure usage in different clinical settings. Moreover, hybrid AI strategies that combine deep learning with expert systems and interpretable models may be helpful in obtaining disclosure and confidence. This final point highlights that the next stage of AI research should revolve around the development of systems that are not only clinically aligned, ethically responsible, and data-efficient, but also make a substantial contribution to medical decision-making and patient safety.

6.5 Final Conclusions

6.5.1 Principal Contributions

This research has several significant implications concerning AI in medicine and the field translation. Primarily, it offers a rigorous comparative analysis of the top CNN architectures under the same experimental conditions, thereby eliminating the confounding factor which has been the major challenge of architectural analyses in the past. The

systematic approach used in this research sets the methodological standards which may raise the quality and comparability of future research.

Secondly, the study demonstrates that current CNN techniques are severely compromised when they are applied to real medical data with an extreme class imbalance. The findings, therefore, challenge the positive evaluation that is typically made in the medical AI literature and provide a valuable background to account for the gap between technological advancement and the capacity to implement it in clinical settings.

Thirdly, the overall assessment system that considers the performance, interpretability, and computational efficiency aspects is a comprehensive method of appraisal, which takes into account the real-world implementation factors besides the strictly technical metrics. With such a scheme as a starting point, the use of scientific evidence for decision-making regarding clinical AI can be facilitated.

6.5.2 Critical Insights

The study discovers that a high overall accuracy can be very misleading in the case of medical AI as it can hide a complete failure of the system in clinically relevant tasks. The consequences of this finding are fundamental to the measuring, reporting, and selection of medical AI systems for clinics.

The lack of success of various architecture solutions in the presence of a harsh class imbalance shows that instructional-methodological breakthroughs are even more important than innovative architecture in realizing clinical usefulness. This observation implies the necessity of radical realignment of research priorities in the direction of dealing with challenges of class imbalance.

These issues of questionable clinical importance of AI attention patterns evoke serious concerns over the compatibility of the current applied methods with clinical

application due to the necessity of decision-making transparency and interpretability to ensure patient safety and responsibility related to professional practice.

6.5.3 Imperative for Change

The results of this research show that the present CNN methods of detecting melanoma are not clinically and need improvement in their basis so that it became possible to ensure patient safety. False negative of up to 99 percent that occurs in all architectures represent unacceptable risks that outweigh possible advantages of better specificity or efficiency.

Further AI medical research should be less focused on sophisticated technical proofs and rather on the elimination of key obstacles to clinical adoption. This extends to giving primacy to addressing the class imbalance, creation of appropriate evaluation models and creation of stringent validation models that give priority to patient safety.

Regulatory bodies, healthcare organisations and funding agencies should be aware of the limits which are found in this study and correct the expectations, requirements and priorities according to the findings in this study. Early introduction of defective AI systems can cause harm to patients and decrease the trust in AI that brings positive outcomes to healthcare and wasted resources.

6.5.4 Hope for Future Progress

Although this paper reports severe drawbacks of the existing methods, it has useful information that may be used to steer future studies to solutions that could be clinically viable. Structured assessment system, the thought, and technical reports provide a mechanism of action toward improving methods of research.

Through identifying certain technical challenges and architectural features, a definite idea about the improvement activity is acquired. Although it is not high enough to

apply in a clinical setting, the relative advantage of EfficientNetB0 illustrates that certain kinds of designs may have the potential to help overcome the issue of class imbalance.

The study, accompanied by a detailed methodology and in-depth analysis, should be a very realistic basis to measure any future steps toward it. By providing an honest account of the existing risks, the study enables the medical AI community to ask the right questions and prioritise their activities in the face of the most difficult barriers to clinical translation.

The ultimate vision of the healthcare system to be improved by AI-assisted diagnosis is still there but quite distant, and it is necessary basic changes in the training of staff, testing of skills and the ways of working with the clinical environment. This research gives the scientific basis that should be used as the correct steps towards that goal while taking into account patient safety and medical usefulness.

6.5.5 Overall Chapter Conclusion

Essentially, this chapter outlines the key points, impact, and proposals of the study arranged in a logical structure that can be used as a roadmap for the next AI in medical diagnostics. Even though the application of CNN-based melanoma detection hardly has any clinical trials that can be trusted, the knowledge gained, nevertheless, points out the key factors for the experimental and ethical practice to be enhanced.

To achieve AI applications suitable for clinical settings, one must combine computational innovations with human expertise, sound data governance, and continuous validation. The key message articulated by the research is that a revolutionary medical AI made in a responsible way cannot be just a technical advancement, but rather a result of ongoing interdisciplinary collaboration; therefore, it guarantees that the integration of AI, as a part of the healthcare systems, is safe, equitable, and beneficial to everyone.

REFERENCES

- Academies, N. *et al.* (2017) *Fostering integrity in research*. National Academies Press.
- Adadi, A. and Berrada, M. (2018) ‘Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI)’, *IEEE Access* [Preprint]. Available at: <https://doi.org/10.1109/ACCESS.2018.2870052>.
- Adamson, A.S. and Smith, A. (2018) ‘Machine Learning and Health Care Disparities in Dermatology’, *JAMA Dermatology*, 154(11), p. 1247. Available at: <https://doi.org/10.1001/jamadermatol.2018.2348>.
- Ahmad, M.A. *et al.* (2021) ‘Artificial Intelligence/Machine Learning (AI/ML)-Based Software as a Medical Device (SaMD) Action Plan’, (January), pp. 4023–4024. Available at: <https://doi.org/10.1145/3447548.3470823>.
- Ahmed, M.I. *et al.* (2023) ‘A Systematic Review of the Barriers to the Implementation of Artificial Intelligence in Healthcare.’, *Cureus*, 15(10), p. e46454. Available at: <https://doi.org/10.7759/cureus.46454>.
- Ahuja, A.S. (2019) ‘The impact of artificial intelligence in medicine on the future role of the physician’, *PeerJ*, 7, p. e7702. Available at: <https://doi.org/10.7717/peerj.7702>.
- Altman, D.G. and Bland, J.M. (1994) ‘Statistics Notes: Diagnostic tests 1: sensitivity and specificity’, *BMJ*, 308(6943), pp. 1552–1552. Available at: <https://doi.org/10.1136/bmj.308.6943.1552>.
- Amann, J. *et al.* (2020) ‘Explainability for artificial intelligence in healthcare: a multidisciplinary perspective’, *BMC Medical Informatics and Decision Making* [Preprint]. Available at: <https://doi.org/10.1186/s12911-020-01332-6>.
- Amodei, D. *et al.* (2016) ‘Concrete Problems in AI Safety’.
- Ancona, M. *et al.* (2017) ‘Towards better understanding of gradient-based attribution methods for deep neural networks’, *arXiv preprint arXiv:1711.06104* [Preprint].
- Anders, C.J. *et al.* (2023) ‘Software for Dataset-wide XAI: From Local Explanations to Global Insights with Zennit, CoRelAy, and ViRelAy’.
- Argenziano, G. *et al.* (2003) ‘Dermoscopy of pigmented skin lesions: Results of a consensus meeting via the Internet’, *Journal of the American Academy of Dermatology*, 48(5), pp. 679–693. Available at: <https://doi.org/10.1067/mjd.2003.281>.
- Assadullah, M.M. (2019) ‘Barriers to Artificial Intelligence Adoption in Healthcare Management: A Systematic Review’, *SSRN Electronic Journal* [Preprint]. Available at: <https://doi.org/10.2139/ssrn.3530598>.
- Azizi, S. *et al.* (2021) ‘Big Self-Supervised Models Advance Medical Image Classification’, in *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*. IEEE, pp. 3458–3468. Available at: <https://doi.org/10.1109/ICCV48922.2021.00346>.
- Balch, C.M. *et al.* (2009) ‘Final Version of 2009 AJCC Melanoma Staging and Classification’, *Journal of Clinical Oncology*, 27(36), pp. 6199–6206. Available at: <https://doi.org/10.1200/JCO.2009.23.4799>.
- Barandela, R. *et al.* (2003) ‘Strategies for learning in class imbalance problems’, *Pattern Recognition*, 36(3), pp. 849–851. Available at: <https://doi.org/10.1016/S0031->

3203(02)00257-1.

- Barata, C. *et al.* (2014) ‘Two Systems for the Detection of Melanomas in Dermoscopy Images Using Texture and Color Features’, *IEEE Systems Journal*, 8(3), pp. 965–979. Available at: <https://doi.org/10.1109/JSYST.2013.2271540>.
- Beam *et al.* (2020) ‘Challenges to the Reproducibility of Machine Learning Models in Health Care’, *JAMA*, 323(4), p. 305. Available at: <https://doi.org/10.1001/jama.2019.20866>.
- Beam, A.L. and Kohane, I.S. (2018) ‘Big Data and Machine Learning in Health Care’, *JAMA*, 319(13), p. 1317. Available at: <https://doi.org/10.1001/jama.2017.18391>.
- Beauchamp, T. and Childress, J. (2019) ‘Principles of biomedical ethics: marking its fortieth anniversary’, *The American journal of bioethics*. Taylor & Francis, pp. 9–12.
- Benavoli, A. *et al.* (2017) ‘Time for a change: a tutorial for comparing multiple classifiers through Bayesian analysis’, *Journal of Machine Learning Research*, 18(77), pp. 1–36.
- Benjamini, Y. and Hochberg, Y. (1995) ‘Controlling the False Discovery Rate: A Practical and Powerful Approach to Multiple Testing’, *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 57(1), pp. 289–300. Available at: <https://doi.org/10.1111/j.2517-6161.1995.tb02031.x>.
- Bianco, S. *et al.* (2018) ‘Benchmark Analysis of Representative Deep Neural Network Architectures’, *IEEE Access*, 6, pp. 64270–64277. Available at: <https://doi.org/10.1109/ACCESS.2018.2877890>.
- Bishop, C.M. (2006) *Pattern recognition and machine learning*. Springer Science+Business Media, LLC Berlin, Germany:
- Bissoto, A. *et al.* (2018) ‘Deep-Learning Ensembles for Skin-Lesion Segmentation, Analysis, Classification: RECOD Titans at ISIC Challenge 2018’.
- Borenstein, M. *et al.* (2009) *Introduction to meta-analysis*. Wiley. Available at: <https://doi.org/10.1002/9780470743386>.
- Bouthillier, X. *et al.* (2021) ‘Accounting for variance in machine learning benchmarks’, *Proceedings of Machine Learning and Systems*, 3, pp. 747–769.
- Box, G.E.P., Hunter, J.S. and Hunter, W.G. (2005) ‘Statistics for experiments: design, innovation, and discovery’, *Hoboken: Wiley* [Preprint].
- Bradley, A.P. (1997) ‘The use of the area under the ROC curve in the evaluation of machine learning algorithms’, *Pattern Recognition* [Preprint]. Available at: [https://doi.org/10.1016/S0031-3203\(96\)00142-2](https://doi.org/10.1016/S0031-3203(96)00142-2).
- Brinker, Titus J. *et al.* (2019) ‘A convolutional neural network trained with dermoscopic images performed on par with 145 dermatologists in a clinical melanoma image classification task’, *European Journal of Cancer*, 111, pp. 148–154. Available at: <https://doi.org/10.1016/j.ejca.2019.02.005>.
- Brinker, Titus J *et al.* (2019) ‘Deep learning outperformed 136 of 157 dermatologists in a head-to-head dermoscopic melanoma image classification task’, *European Journal of Cancer*, 113, pp. 47–54. Available at: <https://doi.org/https://doi.org/10.1016/j.ejca.2019.04.001>.
- Brodersen, K.H. *et al.* (2010) ‘The Balanced Accuracy and Its Posterior Distribution’, in

- 2010 20th International Conference on Pattern Recognition. IEEE, pp. 3121–3124. Available at: <https://doi.org/10.1109/ICPR.2010.764>.
- Bryman, A. (2016) *Social Research Methods*. 5th edn. Oxford, UK: Oxford University Press.
- Buades, A., Coll, B. and Morel, J.-M. (2005) ‘A non-local algorithm for image denoising’, in *2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR’05)*, pp. 60–65 vol. 2. Available at: <https://doi.org/10.1109/CVPR.2005.38>.
- Buda, M., Maki, A. and Mazurowski, M.A. (2018) ‘A systematic study of the class imbalance problem in convolutional neural networks’, *Neural Networks*, 106, pp. 249–259. Available at: <https://doi.org/10.1016/j.neunet.2018.07.011>.
- Buslaev, A. *et al.* (2020) ‘Albumentations: Fast and Flexible Image Augmentations’, *Information*, 11(2), p. 125. Available at: <https://doi.org/10.3390/info11020125>.
- Cabitz, *et al.* (2017) ‘Unintended Consequences of Machine Learning in Medicine’, *JAMA*, 318(6), p. 517. Available at: <https://doi.org/10.1001/jama.2017.7797>.
- Cai, H. *et al.* (2019) ‘Once-for-all: Train one network and specialize it for efficient deployment’, *arXiv preprint arXiv:1908.09791* [Preprint].
- Campbell, D.T. and Stanley, J. (1963) ‘Experimental and quasi-experimental designs for research on teaching’, *Handbook of Research on Teaching*, edited by N.L. Gage, pp. 171–246.
- Canziani, A., Paszke, A. and Culurciello, E. (2017) ‘An Analysis of Deep Neural Network Models for Practical Applications’.
- Carli, P. *et al.* (2004) ‘Addition of dermoscopy to conventional naked-eye examination in melanoma screening: a randomized study’, *Journal of the American Academy of Dermatology*, 50(5), pp. 683–689. Available at: <https://doi.org/10.1016/j.jaad.2003.09.009>.
- Carpenter, J. and Bithell, J. (2000) ‘Bootstrap confidence intervals: when, which, what? A practical guide for medical statisticians’, *Statistics in medicine*, 19(9), pp. 1141–1164.
- Caruana, R., Lawrence, S. and Giles, L. (2001) ‘Overfitting in neural nets: Backpropagation, conjugate gradient, and early stopping’, *Advances in Neural Information Processing Systems* [Preprint].
- Cassidy, B. *et al.* (2022) ‘Analysis of the ISIC image datasets: Usage, benchmarks and recommendations’, *Medical Image Analysis*, 75, p. 102305. Available at: <https://doi.org/10.1016/j.media.2021.102305>.
- Celebi, M.E. *et al.* (2015) ‘A state-of-the-art survey on lesion border detection in dermoscopy images’, *Dermoscopy image analysis*, 10(1), pp. 97–129.
- Chamberlain, A.J. *et al.* (2003) ‘Nodular melanoma: Patients’ perceptions of presenting features and implications for earlier detection’, *Journal of the American Academy of Dermatology*, 48(5), pp. 694–701. Available at: <https://doi.org/10.1067/mjd.2003.216>.
- Char *et al.* (2018) ‘Implementing Machine Learning in Health Care — Addressing Ethical Challenges’, *New England Journal of Medicine*, 378(11), pp. 981–983. Available at: <https://doi.org/10.1056/NEJMp1714229>.

- Chatellier, G. *et al.* (1996) ‘The number needed to treat: a clinically useful nomogram in its proper context’, *BMJ*, 312(7028), pp. 426–429. Available at: <https://doi.org/10.1136/bmj.312.7028.426>.
- Chatfield, C. (2000) *Time-series forecasting*. Chapman and Hall/CRC.
- Chattopadhyay, A. *et al.* (2018) ‘Grad-CAM++: Generalized Gradient-Based Visual Explanations for Deep Convolutional Networks’, in *2018 IEEE Winter Conference on Applications of Computer Vision (WACV)*. IEEE, pp. 839–847. Available at: <https://doi.org/10.1109/WACV.2018.00097>.
- Chawla, N. V. *et al.* (2002) ‘SMOTE: Synthetic Minority Over-sampling Technique’, *Journal of Artificial Intelligence Research*, 16, pp. 321–357. Available at: <https://doi.org/10.1613/jair.953>.
- Chen *et al.* (2020) ‘Treating health disparities with artificial intelligence’, *Nature Medicine*, 26(1), pp. 16–17. Available at: <https://doi.org/10.1038/s41591-019-0649-2>.
- Chen, J.H. and Asch, S.M. (2017) ‘Machine Learning and Prediction in Medicine — Beyond the Peak of Inflated Expectations’, *New England Journal of Medicine*, 376(26), pp. 2507–2509. Available at: <https://doi.org/10.1056/NEJMp1702071>.
- Chen, P.-H.C., Liu, Y. and Peng, L. (2019) ‘How to develop machine learning models for healthcare’, *Nature Materials*, 18(5), pp. 410–414. Available at: <https://doi.org/10.1038/s41563-019-0345-0>.
- Chen, R.J. *et al.* (2021) ‘Synthetic data in machine learning for medicine and healthcare’, *Nature Biomedical Engineering* [Preprint]. Available at: <https://doi.org/10.1038/s41551-021-00751-8>.
- Chetlur, S. *et al.* (2014) ‘cuDNN: Efficient Primitives for Deep Learning’.
- Chicco, D. and Jurman, G. (2020) ‘The advantages of the Matthews correlation coefficient (MCC) over F1 score and accuracy in binary classification evaluation’, *BMC Genomics*, 21(1), p. 6. Available at: <https://doi.org/10.1186/s12864-019-6413-7>.
- Chollet, F. (2017) ‘Xception: Deep Learning with Depthwise Separable Convolutions’, in *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*. IEEE, pp. 1800–1807. Available at: <https://doi.org/10.1109/CVPR.2017.195>.
- Choudhury, A. (2022) ‘Factors influencing clinicians’ willingness to use an AI-based clinical decision support system’, *Frontiers in Digital Health*, 4. Available at: <https://doi.org/10.3389/fdgth.2022.920662>.
- Chronopoulou, A., Baziotis, C. and Potamianos, A. (2019) ‘An Embarrassingly Simple Approach for Transfer Learning from Pretrained Language Models’, in *Proceedings of the 2019 Conference of the North*. Stroudsburg, PA, USA: Association for Computational Linguistics, pp. 2089–2095. Available at: <https://doi.org/10.18653/v1/N19-1213>.
- Codella, N. *et al.* (2019) ‘Skin Lesion Analysis Toward Melanoma Detection 2018: A Challenge Hosted by the International Skin Imaging Collaboration (ISIC)’.
- Codella, N.C.F. *et al.* (2018) ‘Skin lesion analysis toward melanoma detection: A challenge at the 2017 International symposium on biomedical imaging (ISBI), hosted by the international skin imaging collaboration (ISIC)’, in *2018 IEEE 15th International Symposium on Biomedical Imaging (ISBI 2018)*. IEEE, pp. 168–172. Available at:

- <https://doi.org/10.1109/ISBI.2018.8363547>.
- Cohen, J. (1998) ‘Statistical power analysis for the behavioral sciences’, *New York* [Preprint].
- Combalia, M. *et al.* (2019) ‘BCN20000: Dermoscopic Lesions in the Wild’.
- Committee on Publication Ethics (COPE) (2025) ‘COPE Guidance’.
- Cook, R.J. and Sackett, D.L. (1995) ‘The number needed to treat: a clinically useful measure of treatment effect’, *BMJ*, 310(6977), pp. 452–454. Available at: <https://doi.org/10.1136/bmj.310.6977.452>.
- Cook, T. and Campbell, D.T. (1979) ‘Quasi-Experimentation: Design and Analysis Issues for Field Settings’, *Journal of Personality Assessment* [Preprint].
- Cooper, H., Hedges, L. V and Valentine, J.C. (2019) *The handbook of research synthesis and meta-analysis*. Russell Sage Foundation.
- Cooper, H., Hedges, L.V. and Valentine, J.C. (2009) ‘The handbook of research synthesis and meta-analysis’, in *The Hand. of Res. Synthesis and Meta-Analysis, 2nd Ed.* Russell Sage Foundation, pp. 1–615.
- Creswell, J.W. and Creswell, J.D. (2017) *Research design: Qualitative, quantitative, and mixed methods approaches*. Sage publications.
- Creswell, J.W. and Plano Clark, V.L. (2018) *Designing and Conducting Mixed Methods Research*. 3rd edn. SAGE Publications.
- Cronbach, L.J. and Meehl, P.E. (1955) ‘Construct validity in psychological tests.’, *Psychological Bulletin*, 52(4), pp. 281–302. Available at: <https://doi.org/10.1037/h0040957>.
- Cubuk, E.D. *et al.* (2020) ‘Randaugment: Practical automated data augmentation with a reduced search space’, in *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*. IEEE, pp. 3008–3017. Available at: <https://doi.org/10.1109/CVPRW50498.2020.00359>.
- D’Amour, A. *et al.* (2020) ‘Underspecification Presents Challenges for Credibility in Modern Machine Learning’.
- Dabov, K. *et al.* (2007) ‘Image Denoising by Sparse 3-D Transform-Domain Collaborative Filtering’, *IEEE Transactions on Image Processing*, 16(8), pp. 2080–2095. Available at: <https://doi.org/10.1109/TIP.2007.901238>.
- Davis, J. and Goadrich, M. (2006) ‘The relationship between Precision-Recall and ROC curves’, in *Proceedings of the 23rd international conference on Machine learning - ICML ’06*. New York, New York, USA: ACM Press, pp. 233–240. Available at: <https://doi.org/10.1145/1143844.1143874>.
- Davis, S.E. *et al.* (2017) ‘Calibration drift in regression and machine learning models for acute kidney injury’, *Journal of the American Medical Informatics Association*, 24(6), pp. 1052–1061. Available at: <https://doi.org/10.1093/jamia/ocx030>.
- Deb, K. *et al.* (2002) ‘A fast and elitist multiobjective genetic algorithm: NSGA-II’, *IEEE Transactions on Evolutionary Computation*, 6(2), pp. 182–197. Available at: <https://doi.org/10.1109/4235.996017>.
- DeGrave, A.J., Janizek, J.D. and Lee, S.-I. (2021) ‘AI for radiographic COVID-19 detection selects shortcuts over signal’, *Nature Machine Intelligence*, 3(7), pp. 610–619. Available at: <https://doi.org/10.1038/s42256-021-00338-7>.

- Demšar, J. (2006) ‘Statistical comparisons of classifiers over multiple data sets’, *Journal of Machine learning research*, 7(Jan), pp. 1–30.
- Deng, J. *et al.* (2010) ‘ImageNet: A large-scale hierarchical image database’, in. Available at: <https://doi.org/10.1109/cvpr.2009.5206848>.
- Dennis, L.K. *et al.* (2008) ‘Sunburns and risk of cutaneous melanoma: does age matter? A comprehensive meta-analysis’, *Annals of epidemiology*, 18(8), pp. 614–627.
- Doi, K. (2007) ‘Computer-aided diagnosis in medical imaging: Historical review, current status and future potential’, *Computerized Medical Imaging and Graphics*, 31(4–5), pp. 198–211. Available at: <https://doi.org/10.1016/j.compmedimag.2007.02.002>.
- Donahue, J. *et al.* (2014) ‘Decaf: A deep convolutional activation feature for generic visual recognition’, in *International conference on machine learning*, pp. 647–655.
- Doshi-Velez, F. and Kim, B. (2017) ‘Towards A Rigorous Science of Interpretable Machine Learning’, (ML), pp. 1–13. Available at: <http://arxiv.org/abs/1702.08608>.
- Dove, E.S., Knoppers, B.M. and Zawati, M.H. (2014) ‘Towards an ethics safe harbor for global biomedical research’, *Journal of Law and the Biosciences*, 1(1), pp. 3–51.
- Drazen, J.M. *et al.* (2010) ‘Uniform format for disclosure of competing interests in ICMJE journals’, *JAMA*, 303(1), pp. 75–76.
- Drummond, M.F. *et al.* (2015) *Methods for the economic evaluation of health care programmes*. Oxford university press.
- Durand, F. and Dorsey, J. (2002) ‘Fast bilateral filtering for the display of high-dynamic-range images’, *ACM Transactions on Graphics*, 21(3), pp. 257–266. Available at: <https://doi.org/10.1145/566654.566574>.
- Dwan, K. *et al.* (2008) ‘Systematic Review of the Empirical Evidence of Study Publication Bias and Outcome Reporting Bias’, *PLoS ONE*. Edited by N. Siegfried, 3(8), p. e3081. Available at: <https://doi.org/10.1371/journal.pone.0003081>.
- Efron, B. and Tibshirani, R.J. (1994) *An introduction to the bootstrap*. Chapman and Hall/CRC.
- Elmore, J.G. *et al.* (2017) ‘Pathologists’ diagnosis of invasive melanoma and melanocytic proliferations: observer accuracy and reproducibility study’, *BMJ*, p. j2813. Available at: <https://doi.org/10.1136/bmj.j2813>.
- Emanuel, E.J. (2000) ‘What Makes Clinical Research Ethical?’, *JAMA*, 283(20), p. 2701. Available at: <https://doi.org/10.1001/jama.283.20.2701>.
- Emanuel, E.J. *et al.* (2020) ‘Fair Allocation of Scarce Medical Resources in the Time of Covid-19’, *New England Journal of Medicine*, 382(21), pp. 2049–2055. Available at: <https://doi.org/10.1056/NEJMs2005114>.
- Emre Celebi, M. *et al.* (2013) ‘Lesion Border Detection in Dermoscopy Images Using Ensembles of Thresholding Methods’, *Skin Research and Technology*, 19(1). Available at: <https://doi.org/10.1111/j.1600-0846.2012.00636.x>.
- Erdmann, F. *et al.* (2013) ‘International trends in the incidence of malignant melanoma 1953–2008—are recent generations at higher or lower risk?’, *International Journal of Cancer*, 132(2), pp. 385–400. Available at: <https://doi.org/10.1002/ijc.27616>.
- Esteva, A. *et al.* (2017) ‘Dermatologist-level classification of skin cancer with deep neural networks’, *Nature* [Preprint]. Available at: <https://doi.org/10.1038/nature21056>.
- Faul, F. *et al.* (2007) ‘G*Power 3: A flexible statistical power analysis program for the

- social, behavioral, and biomedical sciences', *Behavior Research Methods*, 39(2), pp. 175–191. Available at: <https://doi.org/10.3758/BF03193146>.
- De Fauw, J. *et al.* (2018) 'Clinically applicable deep learning for diagnosis and referral in retinal disease', *Nature Medicine*, 24(9), pp. 1342–1350. Available at: <https://doi.org/10.1038/s41591-018-0107-6>.
- Fawcett, T. (2006) 'An introduction to ROC analysis', *Pattern Recognition Letters*, 27(8), pp. 861–874. Available at: <https://doi.org/10.1016/j.patrec.2005.10.010>.
- Feng, H. *et al.* (2018) 'Comparison of Dermatologist Density Between Urban and Rural Counties in the United States', *JAMA Dermatology*, 154(11), p. 1265. Available at: <https://doi.org/10.1001/jamadermatol.2018.3022>.
- Fidalgo Barata, A., Celebi, E. and Marques, J. (2014) 'Improving Dermoscopy Image Classification Using Color Constancy', *IEEE Journal of Biomedical and Health Informatics*, pp. 1–1. Available at: <https://doi.org/10.1109/JBHI.2014.2336473>.
- Figueroa, R.L. *et al.* (2012) 'Predicting sample size required for classification performance', *BMC Medical Informatics and Decision Making*, 12(1), p. 8. Available at: <https://doi.org/10.1186/1472-6947-12-8>.
- Finlayson, S.G. *et al.* (2021) 'The Clinician and Dataset Shift in Artificial Intelligence', *New England Journal of Medicine*, 385(3), pp. 283–286. Available at: <https://doi.org/10.1056/NEJMc2104626>.
- Finnane, A. *et al.* (2017) 'Teledermatology for the Diagnosis and Management of Skin Cancer', *JAMA Dermatology*, 153(3), p. 319. Available at: <https://doi.org/10.1001/jamadermatol.2016.4361>.
- Fisher, R.A. (1935) 'The design of experiments', *London 4th Ed. 125p* [Preprint].
- Flach, P. (2019) 'Performance Evaluation in Machine Learning: The Good, the Bad, the Ugly, and the Way Forward', *Proceedings of the AAAI Conference on Artificial Intelligence*, 33(01), pp. 9808–9814. Available at: <https://doi.org/10.1609/aaai.v33i01.33019808>.
- Floridi, L. *et al.* (2018) 'AI4People—An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations', *Minds and Machines* [Preprint]. Available at: <https://doi.org/10.1007/s11023-018-9482-5>.
- Fort, S., Hu, H. and Lakshminarayanan, B. (2020) 'Deep Ensembles: A Loss Landscape Perspective'.
- Fukushima, K. (1980) 'Neocognitron: A self-organizing neural network model for a mechanism of pattern recognition unaffected by shift in position', *Biological Cybernetics*, 36(4), pp. 193–202. Available at: <https://doi.org/10.1007/BF00344251>.
- Garbe, C. *et al.* (2016) 'Diagnosis and treatment of melanoma. European consensus-based interdisciplinary guideline – Update 2016', *European Journal of Cancer*, 63, pp. 201–217. Available at: <https://doi.org/10.1016/j.ejca.2016.05.005>.
- García, V., Mollineda, R.A. and Sánchez, J.S. (2009) 'Index of Balanced Accuracy: A Performance Measure for Skewed Class Distributions', in, pp. 441–448. Available at: https://doi.org/10.1007/978-3-642-02172-5_57.
- Garnavi, R. *et al.* (2011) 'Border detection in dermoscopy images using hybrid thresholding on optimized color channels', *Computerized Medical Imaging and Graphics*, 35(2), pp. 105–115.

- Gelman, A. *et al.* (2013) *Bayesian Data Analysis*. 3rd Editio.
- Gelman, A. and Hill, J. (2007) *Data analysis using regression and multilevel/hierarchical models*. Cambridge university press.
- Gessert, N. *et al.* (2020) ‘Skin lesion classification using ensembles of multi-resolution EfficientNets with meta data’, *MethodsX*, 7, p. 100864. Available at: <https://doi.org/10.1016/j.mex.2020.100864>.
- Gibbons, J.D. and Chakraborti, S. (2014) *Nonparametric Statistical Inference*. 5th edn. CRC Press.
- Gibelli, F., Maio, G. and Ricci, G. (2024) ‘Editorial: Healthcare in the age of sapient machines: physician decision-making autonomy faced with artificial intelligence. Ethical, deontological and compensatory aspects’, *Frontiers in Medicine*, 11. Available at: <https://doi.org/10.3389/fmed.2024.1477371>.
- Giger, M.L. (2018) ‘Machine Learning in Medical Imaging’, *Journal of the American College of Radiology*, 15(3), pp. 512–520. Available at: <https://doi.org/10.1016/j.jacr.2017.12.028>.
- Girshick, R. *et al.* (2014) ‘Rich feature hierarchies for accurate object detection and semantic segmentation’, in *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*. Available at: <https://doi.org/10.1109/CVPR.2014.81>.
- Glazer, A.M. *et al.* (2017) ‘Clinical Diagnosis of Skin Cancer: Enhancing Inspection and Early Recognition’, *Dermatologic Clinics*, 35(4), pp. 409–416. Available at: <https://doi.org/10.1016/j.det.2017.06.001>.
- Glenning, J. and Gualtieri, L. (2025) ‘Patient Perspectives on Artificial Intelligence in Medical Imaging’, *Journal of Participatory Medicine*, 17, pp. e67816–e67816. Available at: <https://doi.org/10.2196/67816>.
- Glorot, X. and Bengio, Y. (2010) ‘Understanding the difficulty of training deep feedforward neural networks’, in *Journal of Machine Learning Research*.
- Goodfellow, I. *et al.* (2016) *Deep learning*. MIT press Cambridge.
- Goodman, A. *et al.* (2014) ‘Ten Simple Rules for the Care and Feeding of Scientific Data’, *PLoS Computational Biology*. Edited by P.E. Bourne, 10(4), p. e1003542. Available at: <https://doi.org/10.1371/journal.pcbi.1003542>.
- Goyal, P. *et al.* (2018) ‘Accurate, Large Minibatch SGD: Training ImageNet in 1 Hour’.
- Greenland, S. *et al.* (2016) ‘Statistical tests, P values, confidence intervals, and power: a guide to misinterpretations’, *European Journal of Epidemiology*, 31(4), pp. 337–350. Available at: <https://doi.org/10.1007/s10654-016-0149-3>.
- Greenspan, H., van Ginneken, B. and Summers, R.M. (2016) ‘Guest Editorial Deep Learning in Medical Imaging: Overview and Future Promise of an Exciting New Technique’, *IEEE Transactions on Medical Imaging*, 35(5), pp. 1153–1159. Available at: <https://doi.org/10.1109/TMI.2016.2553401>.
- Gulshan, V. *et al.* (2016) ‘Development and Validation of a Deep Learning Algorithm for Detection of Diabetic Retinopathy in Retinal Fundus Photographs.’, *JAMA*, 316(22), pp. 2402–2410. Available at: <https://doi.org/10.1001/jama.2016.17216>.
- Gundersen, O.E. and Kjensmo, S. (2018) ‘State of the art: Reproducibility in artificial intelligence’, in *Proceedings of the AAAI conference on artificial intelligence*.

- Guy, G.P. *et al.* (2015) ‘Prevalence and Costs of Skin Cancer Treatment in the U.S., 2002–2006 and 2007–2011’, *American Journal of Preventive Medicine*, 48(2), pp. 183–187. Available at: <https://doi.org/10.1016/j.amepre.2014.08.036>.
- Guyatt, G.H. *et al.* (2008) ‘GRADE: an emerging consensus on rating quality of evidence and strength of recommendations’, *BMJ*, 336(7650), pp. 924–926. Available at: <https://doi.org/10.1136/bmj.39489.470347.AD>.
- Habli, I., Lawton, T. and Porter, Z. (2020) ‘Artificial intelligence in health care: accountability and safety’, *Bulletin of the World Health Organization*, 98(4), pp. 251–256. Available at: <https://doi.org/10.2471/BLT.19.237487>.
- Hadalgekar, M.S. and Desai, D.N. (2025) ‘Role of Perceived Ease of Use and Perceived Usefulness in Adoption of Mobile Computing Technology’, *International Journal of Social Science and Human Research*, 08(06). Available at: <https://doi.org/10.47191/ijsshr/v8-i6-58>.
- Haenssle, H.A. *et al.* (2018) ‘Man against machine: diagnostic performance of a deep learning convolutional neural network for dermoscopic melanoma recognition in comparison to 58 dermatologists’, *Annals of Oncology*, 29(8), pp. 1836–1842. Available at: <https://doi.org/10.1093/annonc/mdy166>.
- Hastie, T. *et al.* (2009) ‘The elements of statistical learning: Data Mining, Inference, and Prediction’. Springer series in statistics New-York.
- Hathaway, L. and Gregg, M. (2021) ‘Closing the Intention–Behavior Gap’, *ACSM’S Health & Fitness Journal*, 25(4), pp. 37–39. Available at: <https://doi.org/10.1249/FIT.0000000000000676>.
- He, H. and Garcia, E.A. (2009) ‘Learning from Imbalanced Data’, *IEEE Transactions on Knowledge and Data Engineering*, 21(9), pp. 1263–1284. Available at: <https://doi.org/10.1109/TKDE.2008.239>.
- He, K. *et al.* (2015) ‘Delving Deep into Rectifiers: Surpassing Human-Level Performance on ImageNet Classification’, in *2015 IEEE International Conference on Computer Vision (ICCV)*. IEEE, pp. 1026–1034. Available at: <https://doi.org/10.1109/ICCV.2015.123>.
- He, K. *et al.* (2016) ‘Deep residual learning for image recognition’, in *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*. Available at: <https://doi.org/10.1109/CVPR.2016.90>.
- Helton, J.C. *et al.* (2006) ‘Survey of sampling-based methods for uncertainty and sensitivity analysis’, *Reliability Engineering & System Safety*, 91(10–11), pp. 1175–1209. Available at: <https://doi.org/10.1016/j.ress.2005.11.017>.
- Henderson, P. *et al.* (2018) ‘Deep Reinforcement Learning That Matters’, *Proceedings of the AAAI Conference on Artificial Intelligence*, 32(1). Available at: <https://doi.org/10.1609/aaai.v32i1.11694>.
- Hendrycks, D. and Dietterich, T. (2019) ‘Benchmarking neural network robustness to common corruptions and perturbations’, *arXiv preprint arXiv:1903.12261* [Preprint].
- Hernán, M.A. and Robins, J.M. (2020) ‘Causal inference: what if’, *Causal Inference: What If*, pp. 3–12.
- Hestness, J. *et al.* (2017) ‘Deep Learning Scaling is Predictable, Empirically’.
- Hey, T. *et al.* (2009) *The fourth paradigm: data-intensive scientific discovery*. Microsoft

- research Redmond, WA.
- Hochberg, Y. and Tamhane, A.C. (1987) *Multiple Comparison Procedures*. Wiley (Wiley Series in Probability and Statistics). Available at: <https://doi.org/10.1002/9780470316672>.
- Hollander, M., Wolfe, D.A. and Chicken, E. (2013) *Nonparametric Statistical Methods*. 3rd edn. John Wiley & Sons.
- Holm, S. (1979) 'A simple sequentially rejective multiple test procedure', *Scandinavian journal of statistics*, pp. 65–70.
- Holzinger, A. *et al.* (2017) 'What do we need to build explainable AI systems for the medical domain?'
- Hooker, S. *et al.* (2019) 'A benchmark for interpretability methods in deep neural networks', *Advances in neural information processing systems*, 32.
- Howard, A.G. *et al.* (2017) 'MobileNets: Efficient Convolutional Neural Networks for Mobile Vision Applications'.
- Howard, J. and Ruder, S. (2018) 'Universal Language Model Fine-tuning for Text Classification', in *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. Stroudsburg, PA, USA: Association for Computational Linguistics, pp. 328–339. Available at: <https://doi.org/10.18653/v1/P18-1031>.
- Hu, J., Shen, L. and Sun, G. (2018) 'Squeeze-and-Excitation Networks', in *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*. IEEE, pp. 7132–7141. Available at: <https://doi.org/10.1109/CVPR.2018.00745>.
- Huang, S.-C. *et al.* (2020) 'Fusion of medical imaging and electronic health records using deep learning: a systematic review and implementation guidelines', *npj Digital Medicine*, 3(1), p. 136. Available at: <https://doi.org/10.1038/s41746-020-00341-z>.
- Hunink, M.G.M. *et al.* (2014) *Decision making in health and medicine: integrating evidence and values*. Cambridge university press.
- Hutson, M. (2018) 'Artificial intelligence faces reproducibility crisis', *Science*, 359(6377), pp. 725–726. Available at: <https://doi.org/10.1126/science.359.6377.725>.
- Ignatov, A. *et al.* (2019) 'AI Benchmark: Running Deep Neural Networks on Android Smartphones', in, pp. 288–314. Available at: https://doi.org/10.1007/978-3-030-11021-5_19.
- International Committee of Medical Journal Editors (ICMJE) (2025) 'Recommendations for the Conduct, Reporting, Editing, and Publication of Scholarly Work in Medical Journals'.
- Intersoft consulting (2025) 'General Data Protection Regulation (GDPR) – Legal Text'.
- Ioffe, S. and Szegedy, C. (2015) 'Batch normalization: Accelerating deep network training by reducing internal covariate shift', in *32nd International Conference on Machine Learning, ICML 2015*.
- Jacobsen, A. *et al.* (2020) 'FAIR principles: interpretations and implementation considerations', *Data intelligence*. MIT Press One Rogers Street, Cambridge, MA 02142-1209, USA journals-info~..., pp. 10–29.
- Jacobson, N.S. and Truax, P. (1991) 'Clinical significance: A statistical approach to defining meaningful change in psychotherapy research.', *Journal of Consulting and*

- Clinical Psychology*, 59(1), pp. 12–19. Available at: <https://doi.org/10.1037/0022-006X.59.1.12>.
- Jaimes, N. *et al.* (2012) ‘Clinical and dermoscopic characteristics of amelanotic melanomas that are not of the nodular subtype’, *Journal of the European Academy of Dermatology and Venereology*, 26(5), pp. 591–596. Available at: <https://doi.org/10.1111/j.1468-3083.2011.04122.x>.
- Japkowicz, N. and Stephen, S. (2002) ‘The class imbalance problem: A systematic study’, *Intelligent Data Analysis*, 6(5), pp. 429–449. Available at: <https://doi.org/10.3233/IDA-2002-6504>.
- Jobin, A., Ienca, M. and Vayena, E. (2019) ‘The global landscape of AI ethics guidelines’, *Nature Machine Intelligence* [Preprint]. Available at: <https://doi.org/10.1038/s42256-019-0088-2>.
- Johnson, J.M. and Khoshgoftaar, T.M. (2019) ‘Survey on deep learning with class imbalance’, *Journal of Big Data*, 6(1), p. 27. Available at: <https://doi.org/10.1186/s40537-019-0192-5>.
- Johnson, R.A. and Wichern, D.W. (2007) ‘Applied multivariate statistical analysis’.
- Johnson, R.B. and Onwuegbuzie, A.J. (2004) ‘Mixed Methods Research: A Research Paradigm Whose Time Has Come’, *Educational Researcher*, 33(7), pp. 14–26. Available at: <https://doi.org/10.3102/0013189X033007014>.
- Jouppi, N.P. *et al.* (2017) ‘In-Datacenter Performance Analysis of a Tensor Processing Unit’, in *Proceedings of the 44th Annual International Symposium on Computer Architecture*. New York, NY, USA: ACM, pp. 1–12. Available at: <https://doi.org/10.1145/3079856.3080246>.
- Kaiser, L., Gomez, A.N. and Chollet, F. (2017) ‘Depthwise Separable Convolutions for Neural Machine Translation’.
- Kaissis, G.A. *et al.* (2020) ‘Secure, privacy-preserving and federated machine learning in medical imaging’, *Nature Machine Intelligence*, 2(6), pp. 305–311. Available at: <https://doi.org/10.1038/s42256-020-0186-1>.
- Kapoor, S. and Narayanan, A. (2023) ‘Leakage and the reproducibility crisis in machine-learning-based science’, *Patterns*, 4(9), p. 100804. Available at: <https://doi.org/10.1016/j.patter.2023.100804>.
- Kather, J.N. *et al.* (2019) ‘Deep learning can predict microsatellite instability directly from histology in gastrointestinal cancer’, *Nature Medicine*, 25(7), pp. 1054–1056. Available at: <https://doi.org/10.1038/s41591-019-0462-y>.
- Kattan, M.W. and Vickers, A.J. (2004) ‘Incorporating predictions of individual patient risk in clinical trials’, in *Urologic Oncology: Seminars and Original Investigations*, pp. 348–352.
- Kazdin, A.E. (1999) ‘The meanings and measurement of clinical significance.’, *Journal of Consulting and Clinical Psychology*, 67(3), pp. 332–339. Available at: <https://doi.org/10.1037/0022-006X.67.3.332>.
- Keane *et al.* (2018) ‘With an eye to AI and autonomous diagnosis’, *npj Digital Medicine*, 1(1), p. 40. Available at: <https://doi.org/10.1038/s41746-018-0048-y>.
- Ker, J. *et al.* (2018) ‘Deep Learning Applications in Medical Image Analysis’, *IEEE Access*, 6, pp. 9375–9389. Available at:

- <https://doi.org/10.1109/ACCESS.2017.2788044>.
- Keys, R. (1981) ‘Cubic convolution interpolation for digital image processing’, *IEEE Transactions on Acoustics, Speech, and Signal Processing*, 29(6), pp. 1153–1160. Available at: <https://doi.org/10.1109/TASSP.1981.1163711>.
- Khanna, N.N. *et al.* (2022) ‘Economics of Artificial Intelligence in Healthcare: Diagnosis vs. Treatment’, *Healthcare*, 10(12), p. 2493. Available at: <https://doi.org/10.3390/healthcare10122493>.
- Kimball, A.B. and Resneck, J.S. (2008) ‘The US dermatology workforce: A specialty remains in shortage’, *Journal of the American Academy of Dermatology*, 59(5), pp. 741–745. Available at: <https://doi.org/10.1016/j.jaad.2008.06.037>.
- Kindermans, P.-J. *et al.* (2019) ‘The (Un)reliability of Saliency Methods’, in, pp. 267–280. Available at: https://doi.org/10.1007/978-3-030-28954-6_14.
- Kingma, D.P. and Ba, J. (2017) ‘Adam: A Method for Stochastic Optimization’.
- Kittler, H. *et al.* (2002) ‘Diagnostic accuracy of dermoscopy’, *The Lancet Oncology*, 3(3), pp. 159–165. Available at: [https://doi.org/10.1016/S1470-2045\(02\)00679-4](https://doi.org/10.1016/S1470-2045(02)00679-4).
- Klein, J.P. and Moeschberger, M.L. (2003) *Survival analysis: techniques for censored and truncated data*. New York, NY: Springer New York (Statistics for Biology and Health). Available at: <https://doi.org/10.1007/b97377>.
- Kohavi, R. and others (1995) ‘A study of cross-validation and bootstrap for accuracy estimation and model selection’, in *Ijcai*, pp. 1137–1145.
- Koonce, B. (2021) ‘EfficientNet’, in *Convolutional Neural Networks with Swift for Tensorflow*. Berkeley, CA: Apress, pp. 109–123. Available at: https://doi.org/10.1007/978-1-4842-6168-2_10.
- Krawczyk, B. (2016) ‘Learning from imbalanced data: open challenges and future directions’, *Progress in Artificial Intelligence*, 5(4), pp. 221–232. Available at: <https://doi.org/10.1007/s13748-016-0094-0>.
- Krizhevsky, A., Sutskever, I. and Hinton, G.E. (2012) ‘ImageNet Classification with Deep Convolutional Neural Networks’, *Advances in neural information processing systems*, pp. 1–1432. Available at: <https://doi.org/10.1201/9781420010749>.
- Krstajic, D. *et al.* (2014) ‘Cross-validation pitfalls when selecting and assessing regression and classification models’, *Journal of Cheminformatics*, 6(1), p. 10. Available at: <https://doi.org/10.1186/1758-2946-6-10>.
- Kruschke, J.K. (2014) *Doing Bayesian data analysis: A tutorial with R, JAGS*. 2nd edn. Available at: <https://doi.org/10.1016/B978-0-12-405888-0.09999-2>.
- Kubat, M. and Matwin, S. (1997) ‘Addressing the curse of imbalanced training sets: one-sided selection’, in *Proceedings of the 14th international conference on machine learning*, pp. 179–186.
- Kuhn, T.S. (1962) ‘The Structure of Scientific Revolutions’.
- Lapuschkin, S. *et al.* (2019) ‘Unmasking Clever Hans predictors and assessing what machines really learn’, *Nature Communications*, 10(1), p. 1096. Available at: <https://doi.org/10.1038/s41467-019-08987-4>.
- Larrazabal, A.J. *et al.* (2020) ‘Gender imbalance in medical imaging datasets produces biased classifiers for computer-aided diagnosis’, *Proceedings of the National Academy of Sciences*, 117(23), pp. 12592–12594. Available at:

- <https://doi.org/10.1073/pnas.1919012117>.
- LeCun, Y. *et al.* (1989) 'Backpropagation Applied to Handwritten Zip Code Recognition', *Neural Computation* [Preprint]. Available at: <https://doi.org/10.1162/neco.1989.1.4.541>.
- LeCun, Y., Bengio, Y. and Hinton, G. (2015) 'Deep learning', *Nature*, 521(7553), pp. 436–444. Available at: <https://doi.org/10.1038/nature14539>.
- Lee, A.T., Ramasamy, R.K. and Subbarao, A. (2025) 'Understanding Psychosocial Barriers to Healthcare Technology Adoption: A Review of TAM Technology Acceptance Model and Unified Theory of Acceptance and Use of Technology and UTAUT Frameworks', *Healthcare*, 13(3), p. 250. Available at: <https://doi.org/10.3390/healthcare13030250>.
- Li, T. *et al.* (2020) 'Federated Learning: Challenges, Methods, and Future Directions', *IEEE Signal Processing Magazine* [Preprint]. Available at: <https://doi.org/10.1109/MSP.2020.2975749>.
- Lim, S. *et al.* (2019) 'Fast autoaugment', *Advances in neural information processing systems*, 32.
- Lipton, Z.C. (2018) 'The mythos of model interpretability: In machine learning, the concept of interpretability is both important and slippery.', *Queue* [Preprint].
- Litjens, G. *et al.* (2017) 'A survey on deep learning in medical image analysis', *Medical Image Analysis*, 42, pp. 60–88. Available at: <https://doi.org/10.1016/j.media.2017.07.005>.
- Liu, X. *et al.* (2019) 'A comparison of deep learning performance against health-care professionals in detecting diseases from medical imaging: a systematic review and meta-analysis', *The Lancet Digital Health* [Preprint]. Available at: [https://doi.org/10.1016/S2589-7500\(19\)30123-2](https://doi.org/10.1016/S2589-7500(19)30123-2).
- Liu, X. *et al.* (2020) 'Reporting guidelines for clinical trial reports for interventions involving artificial intelligence: the CONSORT-AI extension', *The Lancet Digital Health*, 2(10), pp. e537–e548. Available at: [https://doi.org/10.1016/S2589-7500\(20\)30218-1](https://doi.org/10.1016/S2589-7500(20)30218-1).
- Loshchilov, I. and Hutter, F. (2019) 'Decoupled Weight Decay Regularization'.
- Luo, W. *et al.* (2016) 'Guidelines for Developing and Reporting Machine Learning Predictive Models in Biomedical Research: A Multidisciplinary View', *Journal of Medical Internet Research*, 18(12), p. e323. Available at: <https://doi.org/10.2196/jmir.5870>.
- Ma, N. *et al.* (2018) 'ShuffleNet V2: Practical Guidelines for Efficient CNN Architecture Design', in, pp. 122–138. Available at: https://doi.org/10.1007/978-3-030-01264-9_8.
- Marchetti, M.A. *et al.* (2018) 'Results of the 2016 International Skin Imaging Collaboration International Symposium on Biomedical Imaging challenge: Comparison of the accuracy of computer algorithms to dermatologists for the diagnosis of melanoma from dermoscopic images', *Journal of the American Academy of Dermatology*, 78(2), pp. 270–277.e1. Available at: <https://doi.org/10.1016/j.jaad.2017.08.016>.
- Masters, K. (2019) 'Artificial intelligence in medical education', *Medical Teacher*, 41(9),

- pp. 976–980. Available at: <https://doi.org/10.1080/0142159X.2019.1595557>.
- Matheny, M., Israni, S.T. and Ahmed, M. (2018) ‘Artificial Intelligence in Health Care: The Hope, the Hype, the Promise, the Peril’, *National Academy of Medicine*, pp. 1–269.
- Matheny, M.E., Whicher, D. and Israni, S.T. (2020) ‘Artificial intelligence in health care: a report from the National Academy of Medicine’, *Jama*, 323(6), pp. 509–510.
- Matthews, B.W. (1975) ‘Comparison of the predicted and observed secondary structure of T4 phage lysozyme’, *Biochimica et Biophysica Acta (BBA) - Protein Structure*, 405(2), pp. 442–451. Available at: [https://doi.org/10.1016/0005-2795\(75\)90109-9](https://doi.org/10.1016/0005-2795(75)90109-9).
- Maxwell, S.E., Delaney, H.D. and Kelley, K. (2017) *Designing experiments and analyzing data: A model comparison perspective*. Routledge.
- McCulloch, W.S. and Pitts, W. (1943) ‘A logical calculus of the ideas immanent in nervous activity’, *The Bulletin of Mathematical Biophysics*, 5(4), pp. 115–133. Available at: <https://doi.org/10.1007/BF02478259>.
- McKinney, S.M. *et al.* (2020) ‘International evaluation of an AI system for breast cancer screening’, *Nature* [Preprint]. Available at: <https://doi.org/10.1038/s41586-019-1799-6>.
- McShane, B.B. *et al.* (2019) ‘Abandon Statistical Significance’, *The American Statistician*, 73(sup1), pp. 235–245. Available at: <https://doi.org/10.1080/00031305.2018.1527253>.
- Meijering, E. (2002) ‘A chronology of interpolation: from ancient astronomy to modern signal and image processing’, *Proceedings of the IEEE*, 90(3), pp. 319–342. Available at: <https://doi.org/10.1109/5.993400>.
- Mendonca, T. *et al.* (2013) ‘PH2 - A dermoscopic image database for research and benchmarking’, in *2013 35th Annual International Conference of the IEEE Engineering in Medicine and Biology Society (EMBC)*. IEEE, pp. 5437–5440. Available at: <https://doi.org/10.1109/EMBC.2013.6610779>.
- Menzies, S. *et al.* (1996) ‘Frequency and Morphologic Characteristics of Invasive Melanomas Lacking Specific Surface Microscopic Features’, *Archives of Dermatology*, 132(10), p. 1178. Available at: <https://doi.org/10.1001/archderm.1996.03890340038007>.
- Messick, S. (1995) ‘Validity of psychological assessment: Validation of inferences from persons’ responses and performances as scientific inquiry into score meaning.’, *American Psychologist*, 50(9), pp. 741–749. Available at: <https://doi.org/10.1037/0003-066X.50.9.741>.
- Miller, R.G. (1981) *Simultaneous Statistical Inference*. New York, NY: Springer New York (Springer Series in Statistics). Available at: <https://doi.org/10.1007/978-1-4613-8122-8>.
- Mongan *et al.* (2020) ‘Checklist for Artificial Intelligence in Medical Imaging (CLAIM): A Guide for Authors and Reviewers’, *Radiology: Artificial Intelligence*, 2(2), p. e200029. Available at: <https://doi.org/10.1148/ryai.2020200029>.
- Montavon, G., Samek, W. and Müller, K.-R. (2018) ‘Methods for interpreting and understanding deep neural networks’, *Digital Signal Processing*, 73, pp. 1–15. Available at: <https://doi.org/10.1016/j.dsp.2017.10.011>.

- Montgomery, D.C. (2017) *Design and analysis of experiments*. John Wiley & sons.
- Montgomery, D.C., Jennings, C.L. and Kulahci, M. (2015) *Introduction to time series analysis and forecasting*. John Wiley & Sons.
- Muller, S.G. and Hutter, F. (2021) ‘TrivialAugment: Tuning-free Yet State-of-the-Art Data Augmentation’, in *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*. IEEE, pp. 754–762. Available at: <https://doi.org/10.1109/ICCV48922.2021.00081>.
- Murdoch, W.J. *et al.* (2019) ‘Definitions, methods, and applications in interpretable machine learning’, *Proceedings of the National Academy of Sciences*, 116(44), pp. 22071–22080. Available at: <https://doi.org/10.1073/pnas.1900654116>.
- Nagendran, M. *et al.* (2020) ‘Artificial intelligence versus clinicians: systematic review of design, reporting standards, and claims of deep learning studies’, *BMJ*, p. m689. Available at: <https://doi.org/10.1136/bmj.m689>.
- National Science Foundation (NSF) (2025) ‘Responsible and Ethical Conduct of Research’.
- Neyman, J. and Pearson, E.S. (1933) ‘On the problem of the most efficient tests of statistical hypotheses’, *Philosophical Transactions of the Royal Society of London. Series A, Containing Papers of a Mathematical or Physical Character*, 231(694–706), pp. 289–337.
- Northcutt, C., Jiang, L. and Chuang, I. (2021) ‘Confident Learning: Estimating Uncertainty in Dataset Labels’, *Journal of Artificial Intelligence Research*, 70, pp. 1373–1411. Available at: <https://doi.org/10.1613/jair.1.12125>.
- Nosek, B.A. *et al.* (2015) ‘Promoting an open research culture’, *Science*, 348(6242), pp. 1422–1425. Available at: <https://doi.org/10.1126/science.aab2374>.
- Nosek, B.A. *et al.* (2018) ‘The preregistration revolution’, *Proceedings of the National Academy of Sciences*, 115(11), pp. 2600–2606. Available at: <https://doi.org/10.1073/pnas.1708274114>.
- O’Hagan, A. *et al.* (2006) ‘Uncertain judgements: eliciting experts’ probabilities’.
- Oakden-Rayner, L. *et al.* (2020) ‘Hidden stratification causes clinically meaningful failures in machine learning for medical imaging’, in *Proceedings of the ACM Conference on Health, Inference, and Learning*. New York, NY, USA: ACM, pp. 151–159. Available at: <https://doi.org/10.1145/3368555.3384468>.
- Obermeyer, Z. and Emanuel, E.J. (2016) ‘Predicting the Future — Big Data, Machine Learning, and Clinical Medicine’, *New England Journal of Medicine*, 375(13), pp. 1216–1219. Available at: <https://doi.org/10.1056/NEJMp1606181>.
- of Intramural Research, N.I. of H. (NIH) (2025) ‘Guidelines for the Conduct of Research Involving Human Subjects’.
- Office for Human Research Protections (OHRP), U.S.D. of H. & H.S. (2025) ‘Federal Policy for the Protection of Human Subjects’.
- Pacheco, A.G.C. *et al.* (2020) ‘PAD-UFES-20: A skin lesion dataset composed of patient data and clinical images collected from smartphones’, *Data in Brief*, 32, p. 106221. Available at: <https://doi.org/10.1016/j.dib.2020.106221>.
- Parikh, R. *et al.* (2008) ‘Understanding and using sensitivity, specificity and predictive values’, *Indian Journal of Ophthalmology*, 56(1), p. 45. Available at:

- <https://doi.org/10.4103/0301-4738.37595>.
- Pascanu, R., Mikolov, T. and Bengio, Y. (2013) ‘On the difficulty of training recurrent neural networks’, *Phylogenetic Diversity: Applications and Challenges in Biodiversity Science*, (2), pp. 41–71.
- Patterson, D. *et al.* (2021) ‘Carbon Emissions and Large Neural Network Training’.
- Peng, R.D. (2011) ‘Reproducible Research in Computational Science’, *Science*, 334(6060), pp. 1226–1227. Available at: <https://doi.org/10.1126/science.1213847>.
- Perez, F. *et al.* (2018) ‘Data Augmentation for Skin Lesion Analysis’, in, pp. 303–311. Available at: https://doi.org/10.1007/978-3-030-01201-4_33.
- Peters, M.E. *et al.* (2018) ‘Deep contextualized word representations’, in *NAACL HLT 2018 - 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies - Proceedings of the Conference*. Available at: <https://doi.org/10.18653/v1/n18-1202>.
- Phillips, M. *et al.* (2019) ‘Assessment of Accuracy of an Artificial Intelligence Algorithm to Detect Melanoma in Images of Skin Lesions’, *JAMA Network Open*, 2(10), p. e1913436. Available at: <https://doi.org/10.1001/jamanetworkopen.2019.13436>.
- Pianosi, F. *et al.* (2016) ‘Sensitivity analysis of environmental models: A systematic review with practical workflow’, *Environmental Modelling & Software*, 79, pp. 214–232. Available at: <https://doi.org/10.1016/j.envsoft.2016.02.008>.
- Pinheiro, J.C. and Bates, D.M. (2000) *Mixed-effects models in S and S-PLUS*. Springer.
- Pizer, S.M. *et al.* (1987) ‘Adaptive histogram equalization and its variations’, *Computer Vision, Graphics, and Image Processing*, 39(3), pp. 355–368. Available at: [https://doi.org/10.1016/S0734-189X\(87\)80186-X](https://doi.org/10.1016/S0734-189X(87)80186-X).
- Pizzichetta, M.A. *et al.* (2004) ‘Amelanotic/hypomelanotic melanoma: clinical and dermoscopic features’, *British Journal of Dermatology*, 150(6), pp. 1117–1124. Available at: <https://doi.org/10.1111/j.1365-2133.2004.05928.x>.
- Popper, K.R. (2002) ‘The logic of scientific discovery’. London: Routledge).(Orig. published as *Logik der Forschung*, Vienna: Julius~....
- Powers, D.M.W. (2020) ‘Evaluation: from precision, recall and F-measure to ROC, informedness, markedness and correlation’, (January 2011). Available at: <https://doi.org/10.9735/2229-3981>.
- Prechelt, L. (1998) ‘Early Stopping - But When?’, in, pp. 55–69. Available at: https://doi.org/10.1007/3-540-49430-8_3.
- Price, W.N., Gerke, S. and Cohen, I.G. (2019) ‘Potential Liability for Physicians Using Artificial Intelligence’, *JAMA*, 322(18), p. 1765. Available at: <https://doi.org/10.1001/jama.2019.15064>.
- Qingfu Zhang and Hui Li (2007) ‘MOEA/D: A Multiobjective Evolutionary Algorithm Based on Decomposition’, *IEEE Transactions on Evolutionary Computation*, 11(6), pp. 712–731. Available at: <https://doi.org/10.1109/TEVC.2007.892759>.
- Raghu, M. *et al.* (2019) ‘Transfusion: Understanding Transfer Learning for Medical Imaging’, in H. Wallach *et al.* (eds) *Advances in Neural Information Processing Systems*. Curran Associates, Inc.
- Rajkomar, A. *et al.* (2018) ‘Ensuring Fairness in Machine Learning to Advance Health Equity’, *Annals of Internal Medicine*, 169(12), pp. 866–872. Available at:

- <https://doi.org/10.7326/M18-1990>.
- Rajpara, S.M. *et al.* (2009) ‘Systematic review of dermoscopy and digital dermoscopy/artificial intelligence for the diagnosis of melanoma’, *British Journal of Dermatology*, 161(3), pp. 591–604. Available at: <https://doi.org/10.1111/j.1365-2133.2009.09093.x>.
- Rajpurkar, P. *et al.* (2017) ‘CheXNet: Radiologist-Level Pneumonia Detection on Chest X-Rays with Deep Learning’.
- Rajpurkar, P. *et al.* (2022) ‘AI in health and medicine’, *Nature Medicine*, 28(1), pp. 31–38. Available at: <https://doi.org/10.1038/s41591-021-01614-0>.
- Ram, K. (2013) ‘Git can facilitate greater reproducibility and increased transparency in science’, *Source Code for Biology and Medicine*, 8(1), p. 7. Available at: <https://doi.org/10.1186/1751-0473-8-7>.
- Rastrelli, M. *et al.* (2014) ‘Melanoma: epidemiology, risk factors, pathogenesis, diagnosis and classification’, *In vivo*, 28(6), pp. 1005–1011.
- Recht, B. *et al.* (2019) ‘Do imagenet classifiers generalize to imagenet?’, in *International conference on machine learning*, pp. 5389–5400.
- Reddi, S.J., Kale, S. and Kumar, S. (2019) ‘On the convergence of adam and beyond’, *arXiv preprint arXiv:1904.09237* [Preprint].
- Reddy, S. *et al.* (2020) ‘A governance model for the application of AI in health care’, *Journal of the American Medical Informatics Association*, 27(3), pp. 491–497. Available at: <https://doi.org/10.1093/jamia/ocz192>.
- Reid, R.J. *et al.* (2024) ‘Actioning the Learning Health System: An applied framework for integrating research into health systems’, *SSM - Health Systems*, 2, p. 100010. Available at: <https://doi.org/10.1016/j.ssmhs.2024.100010>.
- Reinhard, E. *et al.* (2002) ‘Photographic tone reproduction for digital images’, *ACM Transactions on Graphics*, 21(3), pp. 267–276. Available at: <https://doi.org/10.1145/566654.566575>.
- Rencher, A.C. and Christensen, W.F. (2012) *Methods of Multivariate Analysis*. Wiley (Wiley Series in Probability and Statistics). Available at: <https://doi.org/10.1002/9781118391686>.
- Roberts, M. *et al.* (2021) ‘Common pitfalls and recommendations for using machine learning to detect and prognosticate for COVID-19 using chest radiographs and CT scans’, *Nature Machine Intelligence*, 3(3), pp. 199–217. Available at: <https://doi.org/10.1038/s42256-021-00307-0>.
- Rodríguez-Ruiz, A. *et al.* (2019) ‘Detection of Breast Cancer with Mammography: Effect of an Artificial Intelligence Support System’, *Radiology*, 290(2), pp. 305–314. Available at: <https://doi.org/10.1148/radiol.2018181371>.
- Ronneberger, O., Fischer, P. and Brox, T. (2015) ‘U-Net: Convolutional Networks for Biomedical Image Segmentation’, in, pp. 234–241. Available at: https://doi.org/10.1007/978-3-319-24574-4_28.
- Rosenblatt, F. (1958) ‘The perceptron: A probabilistic model for information storage and organization in the brain.’, *Psychological Review*, 65(6), pp. 386–408. Available at: <https://doi.org/10.1037/h0042519>.
- Rotemberg, V. *et al.* (2021) ‘A patient-centric dataset of images and metadata for

- identifying melanomas using clinical context', *Scientific Data*, 8(1), p. 34. Available at: <https://doi.org/10.1038/s41597-021-00815-z>.
- Rothwell, P.M. (2005) 'External validity of randomised controlled trials: "To whom do the results of this trial apply?"', *The Lancet*, 365(9453), pp. 82–93. Available at: [https://doi.org/10.1016/S0140-6736\(04\)17670-8](https://doi.org/10.1016/S0140-6736(04)17670-8).
- Rudin, C. (2019) 'Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead', *Nature Machine Intelligence* [Preprint]. Available at: <https://doi.org/10.1038/s42256-019-0048-x>.
- Russakovsky, O. *et al.* (2015) 'ImageNet Large Scale Visual Recognition Challenge', *International Journal of Computer Vision* [Preprint]. Available at: <https://doi.org/10.1007/s11263-015-0816-y>.
- Saito, T. and Rehmsmeier, M. (2015) 'The Precision-Recall Plot Is More Informative than the ROC Plot When Evaluating Binary Classifiers on Imbalanced Datasets', *PLOS ONE*. Edited by G. Brock, 10(3), p. e0118432. Available at: <https://doi.org/10.1371/journal.pone.0118432>.
- Saltelli, A. *et al.* (2008) *Global sensitivity analysis: the primer*. John Wiley & Sons.
- Samek, W. and Müller, K.-R. (2019) 'Towards Explainable Artificial Intelligence', in, pp. 5–22. Available at: https://doi.org/10.1007/978-3-030-28954-6_1.
- Samek, W., Wiegand, T. and Müller, K.-R. (2017) 'Explainable Artificial Intelligence: Understanding, Visualizing and Interpreting Deep Learning Models'.
- Sanders, G.D. *et al.* (2016) 'Recommendations for Conduct, Methodological Practices, and Reporting of Cost-effectiveness Analyses', *JAMA*, 316(10), p. 1093. Available at: <https://doi.org/10.1001/jama.2016.12195>.
- Sandler, M. *et al.* (2018) 'MobileNetV2: Inverted Residuals and Linear Bottlenecks', in *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*. IEEE, pp. 4510–4520. Available at: <https://doi.org/10.1109/CVPR.2018.00474>.
- Sandve, G.K. *et al.* (2013) 'Ten Simple Rules for Reproducible Computational Research', *PLoS Computational Biology*. Edited by P.E. Bourne, 9(10), p. e1003285. Available at: <https://doi.org/10.1371/journal.pcbi.1003285>.
- Santurkar, S. *et al.* (2018) 'How does batch normalization help optimization?', *Advances in neural information processing systems*, 31.
- Sasaki, Y. (2007) 'The truth of the F-measure', *Teach Tutor Mater* [Preprint].
- Saunders, M., Lewis, P. and Thornhill, A. (2009) *Research methods for business students. Fifth Edition*, Pearson Education, UK.
- Schorr, A. (2023) 'The Technology Acceptance Model (TAM) and its Importance for Digitalization Research: A Review', in *International Symposium on Teknikpsychologie (TecPsy) 2023*. Available at: <https://doi.org/10.2478/9788366675896-005>.
- Schwartz, R. *et al.* (2020) 'Green AI', *Communications of the ACM*, 63(12), pp. 54–63. Available at: <https://doi.org/10.1145/3381831>.
- Scobie, S. and Castle-Clarke, S. (2020) 'Implementing learning health systems in the UK NHS: Policy actions to improve collaboration and transparency and support innovation and better use of analytics', *Learning Health Systems*, 4(1). Available at: <https://doi.org/10.1002/lrh2.10209>.

- Selvaraju, R.R. *et al.* (2017) ‘Grad-CAM: Visual Explanations from Deep Networks via Gradient-Based Localization’, in *2017 IEEE International Conference on Computer Vision (ICCV)*. IEEE, pp. 618–626. Available at: <https://doi.org/10.1109/ICCV.2017.74>.
- Sendak, M.P. *et al.* (2020) ‘Presenting machine learning model information to clinical end users with model facts labels’, *npj Digital Medicine*, 3(1), p. 41. Available at: <https://doi.org/10.1038/s41746-020-0253-3>.
- Seyyed-Kalantari, L. *et al.* (2021) ‘Underdiagnosis bias of artificial intelligence algorithms applied to chest radiographs in under-served patient populations’, *Nature Medicine*, 27(12), pp. 2176–2182. Available at: <https://doi.org/10.1038/s41591-021-01595-0>.
- Shadish, W.R., Cook, T.D. and Campbell, D.T. (2002) *Experimental and quasi-experimental designs for generalized causal inference*. Houghton, Mifflin and Company.
- Shah, N.H., Milstein, A. and Bagley, S.C. (2019) ‘Making Machine Learning Models Clinically Useful’, *JAMA*, 322(14), p. 1351. Available at: <https://doi.org/10.1001/jama.2019.10306>.
- Shah, P. *et al.* (2019) ‘Artificial intelligence and machine learning in clinical development: a translational perspective’, *npj Digital Medicine*, 2(1), p. 69. Available at: <https://doi.org/10.1038/s41746-019-0148-3>.
- Shen, D., Wu, G. and Suk, H.-I. (2017) ‘Deep Learning in Medical Image Analysis’, *Annual Review of Biomedical Engineering*, 19(1), pp. 221–248. Available at: <https://doi.org/10.1146/annurev-bioeng-071516-044442>.
- Shorten, C. and Khoshgoftaar, T.M. (2019) ‘A survey on Image Data Augmentation for Deep Learning’, *Journal of Big Data*, 6(1). Available at: <https://doi.org/10.1186/s40537-019-0197-0>.
- Siegel, R.L. *et al.* (2023a) ‘Cancer statistics, 2023’, *CA: A Cancer Journal for Clinicians*, 73(1), pp. 17–48. Available at: <https://doi.org/10.3322/caac.21763>.
- Siegel, R.L. *et al.* (2023b) ‘Cancer statistics’, *CA: A Cancer Journal for Clinicians*, 73(1), pp. 17–48. Available at: <https://doi.org/10.3322/caac.21763>.
- Simard, P.Y., Steinkraus, D. and Platt, J.C. (2003) ‘Best practices for convolutional neural networks applied to visual document analysis’, in *Seventh International Conference on Document Analysis and Recognition, Proceedings*. IEEE Comput. Soc, pp. 958–963. Available at: <https://doi.org/10.1109/ICDAR.2003.1227801>.
- Simmons, J.P., Nelson, L.D. and Simonsohn, U. (2011) ‘False-Positive Psychology’, *Psychological Science*, 22(11), pp. 1359–1366. Available at: <https://doi.org/10.1177/0956797611417632>.
- Singh, K. *et al.* (2021) ‘Evaluating a Widely Implemented Proprietary Deterioration Index Model among Hospitalized Patients with COVID-19’, *Annals of the American Thoracic Society*, 18(7), pp. 1129–1137. Available at: <https://doi.org/10.1513/AnnalsATS.202006-698OC>.
- Smith, S.L. *et al.* (2018) ‘Don’t Decay the Learning Rate, Increase the Batch Size’.
- Sobaih, A.E.E. *et al.* (2025) ‘Unlocking Patient Resistance to AI in Healthcare: A Psychological Exploration’, *European Journal of Investigation in Health, Psychology and Education*, 15(1), p. 6. Available at:

- <https://doi.org/10.3390/ejihpe15010006>.
- Society, A.C. (2023) ‘Cancer facts & figures’. American Cancer Society Atlanta, GA.
- Song, C. *et al.* (2023) ‘Factors Influencing Instructors’ Adoption and Continued Use of Computing Science Technologies: A Case Study in the Context of Cell Collective’, *CBE—Life Sciences Education*. Edited by G.E. Gardner, 22(3). Available at: <https://doi.org/10.1187/cbe.22-11-0239>.
- Stark, J.A. (2000) ‘Adaptive image contrast enhancement using generalizations of histogram equalization’, *IEEE Transactions on Image Processing*, 9(5), pp. 889–896. Available at: <https://doi.org/10.1109/83.841534>.
- Stodden, V., Seiler, J. and Ma, Z. (2018) ‘An empirical analysis of journal policy effectiveness for computational reproducibility’, *Proceedings of the National Academy of Sciences*, 115(11), pp. 2584–2589. Available at: <https://doi.org/10.1073/pnas.1708290115>.
- Stone, M. (1974) ‘Cross-Validatory Choice and Assessment of Statistical Predictions’, *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 36(2), pp. 111–133. Available at: <https://doi.org/10.1111/j.2517-6161.1974.tb00994.x>.
- Strubell, E., Ganesh, A. and McCallum, A. (2019) ‘Energy and Policy Considerations for Deep Learning in NLP’, in *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*. Stroudsburg, PA, USA: Association for Computational Linguistics, pp. 3645–3650. Available at: <https://doi.org/10.18653/v1/P19-1355>.
- Subbaswamy, A. and Saria, S. (2020) ‘From development to deployment: dataset shift, causality, and shift-stable models in health AI’, *Biostatistics*, 21(2), pp. 345–352.
- Sullivan, G.M. and Feinn, R. (2012) ‘Using Effect Size—or Why the P Value Is Not Enough’, *Journal of Graduate Medical Education*, 4(3), pp. 279–282. Available at: <https://doi.org/10.4300/JGME-D-12-00156.1>.
- Tan, M. and Le, Q. V. (2019) ‘EfficientNet: Rethinking model scaling for convolutional neural networks’, in *36th International Conference on Machine Learning, ICML 2019*.
- Tashakkori, A. and Teddlie, C. (2015) *SAGE Handbook of Mixed Methods in Social & Behavioral Research*, *SAGE Handbook of Mixed Methods in Social & Behavioral Research*. Available at: <https://doi.org/10.4135/9781506335193>.
- Technology, N.I. of S. and (2020) *Security and Privacy Controls for Information Systems and Organizations*. Gaithersburg, MD. Available at: <https://doi.org/10.6028/NIST.SP.800-53r5>.
- Therneau, T.M. and Grambsch, P.M. (2000) *Modeling Survival Data: Extending the Cox Model*. New York, NY: Springer New York (Statistics for Biology and Health). Available at: <https://doi.org/10.1007/978-1-4757-3294-8>.
- Thomas, L. *et al.* (1998) ‘Semiological Value of ABCDE Criteria in the Diagnosis of Cutaneous Pigmented Tumors’, *Dermatology*, 197(1), pp. 11–17. Available at: <https://doi.org/10.1159/000017969>.
- Thompson, N. *et al.* (2023) ‘The Computational Limits of Deep Learning’, in. Available at: <https://doi.org/10.21428/bf6fb269.1f033948>.
- Thomsen, K. *et al.* (2020) ‘Systematic review of machine learning for diagnosis and

- prognosis in dermatology', *Journal of Dermatological Treatment*, 31(5), pp. 496–510. Available at: <https://doi.org/10.1080/09546634.2019.1682500>.
- Tjoa, E. and Guan, C. (2021) 'A Survey on Explainable Artificial Intelligence (XAI): Toward Medical XAI', *IEEE Transactions on Neural Networks and Learning Systems*, 32(11), pp. 4793–4813. Available at: <https://doi.org/10.1109/TNNLS.2020.3027314>.
- Topol, E.J. (2019) 'High-performance medicine: the convergence of human and artificial intelligence', *Nature Medicine*, 25(1), pp. 44–56. Available at: <https://doi.org/10.1038/s41591-018-0300-7>.
- Tschandl, P. *et al.* (2019) 'Comparison of the accuracy of human readers versus machine-learning algorithms for pigmented skin lesion classification: an open, web-based, international, diagnostic study', *The Lancet Oncology*, 20(7), pp. 938–947. Available at: [https://doi.org/10.1016/S1470-2045\(19\)30333-X](https://doi.org/10.1016/S1470-2045(19)30333-X).
- Tschandl, P. *et al.* (2020) 'Human–computer collaboration for skin cancer recognition', *Nature Medicine*, 26(8), pp. 1229–1234. Available at: <https://doi.org/10.1038/s41591-020-0942-0>.
- Tschandl, P., Rosendahl, C. and Kittler, H. (2018) 'The HAM10000 dataset, a large collection of multi-source dermoscopic images of common pigmented skin lesions', *Scientific Data*, 5(1), p. 180161. Available at: <https://doi.org/10.1038/sdata.2018.161>.
- Tun, H.M. *et al.* (2025) 'Trust in Artificial Intelligence–Based Clinical Decision Support Systems Among Health Care Workers: Systematic Review', *Journal of Medical Internet Research*, 27, pp. e69678–e69678. Available at: <https://doi.org/10.2196/69678>.
- U.S. Department of Health and Human Services, Office for Civil Rights (2025) 'Health Information Privacy — HIPAA'.
- Udrea, A. *et al.* (2020) 'Accuracy of a smartphone application for triage of skin lesions based on machine learning algorithms', *Journal of the European Academy of Dermatology and Venereology*, 34(3), pp. 648–655. Available at: <https://doi.org/10.1111/jdv.15935>.
- Varoquaux, G. and Cheplygina, V. (2022) 'Machine learning for medical imaging: methodological failures and recommendations for the future', *npj Digital Medicine*, 5(1), p. 48. Available at: <https://doi.org/10.1038/s41746-022-00592-y>.
- Veit, A., Wilber, M.J. and Belongie, S. (2016) 'Residual networks behave like ensembles of relatively shallow networks', *Advances in neural information processing systems*, 29.
- Verma, K.K. *et al.* (2025) 'Artificial intelligence in melanoma diagnosis: ethical considerations and clinical implementation', *Baylor University Medical Center Proceedings*, 38(4), pp. 577–578. Available at: <https://doi.org/10.1080/08998280.2025.2489873>.
- Wang, Q. *et al.* (2020) 'ECA-Net: Efficient Channel Attention for Deep Convolutional Neural Networks', in *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. IEEE, pp. 11531–11539. Available at: <https://doi.org/10.1109/CVPR42600.2020.01155>.

- Wasserstein, R.L. and Lazar, N.A. (2016) ‘The ASA Statement on p -Values: Context, Process, and Purpose’, *The American Statistician*, 70(2), pp. 129–133. Available at: <https://doi.org/10.1080/00031305.2016.1154108>.
- Whiteman, D.C., Green, A.C. and Olsen, C.M. (2016) ‘The Growing Burden of Invasive Melanoma: Projections of Incidence Rates and Numbers of New Cases in Six Susceptible Populations through 2031’, *Journal of Investigative Dermatology*, 136(6), pp. 1161–1171. Available at: <https://doi.org/10.1016/j.jid.2016.01.035>.
- Wiens, J. *et al.* (2019) ‘Do no harm: a roadmap for responsible machine learning for health care’, *Nature Medicine*, 25(9), pp. 1337–1340. Available at: <https://doi.org/10.1038/s41591-019-0548-6>.
- Wilkinson, M.D. *et al.* (2016) ‘The FAIR Guiding Principles for scientific data management and stewardship’, *Scientific Data*, 3(1), p. 160018. Available at: <https://doi.org/10.1038/sdata.2016.18>.
- Wilson, G. *et al.* (2017) ‘Good enough practices in scientific computing’, *PLOS Computational Biology*. Edited by F. Ouellette, 13(6), p. e1005510. Available at: <https://doi.org/10.1371/journal.pcbi.1005510>.
- Winkler, J.K. *et al.* (2019) ‘Association Between Surgical Skin Markings in Dermoscopic Images and Diagnostic Performance of a Deep Learning Convolutional Neural Network for Melanoma Recognition’, *JAMA Dermatology*, 155(10), p. 1135. Available at: <https://doi.org/10.1001/jamadermatol.2019.1735>.
- Woo, S. *et al.* (2018) ‘CBAM: Convolutional block attention module’, in *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*. Available at: https://doi.org/10.1007/978-3-030-01234-2_1.
- Wu, Z., Shen, C. and van den Hengel, A. (2019) ‘Wider or Deeper: Revisiting the ResNet Model for Visual Recognition’, *Pattern Recognition*, 90, pp. 119–133. Available at: <https://doi.org/10.1016/j.patcog.2019.01.006>.
- Yosinski, J. *et al.* (2014) ‘How transferable are features in deep neural networks?’, in *Advances in Neural Information Processing Systems*.
- You, Y., Gitman, I. and Ginsburg, B. (2017) ‘Large Batch Training of Convolutional Networks’.
- Yu, K.-H. *et al.* (2018) ‘Artificial intelligence in healthcare’, *Nature Biomedical Engineering*, 2(10), pp. 719–731. Available at: <https://doi.org/10.1038/s41551-018-0305-z>.
- Zagoruyko, S. and Komodakis, N. (2017) ‘Wide Residual Networks’.
- Zech, J.R. *et al.* (2018) ‘Variable generalization performance of a deep learning model to detect pneumonia in chest radiographs: A cross-sectional study’, *PLOS Medicine*. Edited by A. Sheikh, 15(11), p. e1002683. Available at: <https://doi.org/10.1371/journal.pmed.1002683>.
- Zhang, C. *et al.* (2017) ‘Understanding deep learning requires rethinking generalization’.
- Zhang, C. *et al.* (2021) ‘Understanding deep learning (still) requires rethinking generalization’, *Communications of the ACM*, 64(3), pp. 107–115. Available at: <https://doi.org/10.1145/3446776>.
- Zhang, H. *et al.* (2019) ‘Theoretically principled trade-off between robustness and

- accuracy', in *International conference on machine learning*, pp. 7472–7482.
- Zhang, J. *et al.* (2018) 'Top-Down Neural Attention by Excitation Backprop', *International Journal of Computer Vision*, 126(10), pp. 1084–1102. Available at: <https://doi.org/10.1007/s11263-017-1059-x>.
- Zhou, H.-Y. *et al.* (2020) 'Comparing to learn: Surpassing imagenet pretraining on radiographs by comparing image representations', in *International Conference on Medical Image Computing and Computer-Assisted Intervention*, pp. 398–407.
- Zuiderveld, K. (1994) 'Contrast Limited Adaptive Histogram Equalization', in *Graphics Gems*. Elsevier, pp. 474–485. Available at: <https://doi.org/10.1016/B978-0-12-336156-1.50061-6>.

APPENDIX A:

A.1 Detailed Statistical Analysis

A.1.1 Bootstrap Confidence Interval Calculations

Table A.1: Complete Bootstrap Analysis Results (1000 iterations)

Architecture	Metric	Mean	Std Dev	95% CI Lower	95% CI Upper	Skewness	Kurtosis
EfficientNetB0	Sensitivity	0.0094	0.0051	0.0043	0.0145	2.34	8.91
EfficientNetB0	Specificity	0.9975	0.0008	0.9967	0.9983	-1.87	5.23
EfficientNetB0	Precision	0.0588	0.0312	0.0276	0.0900	1.98	6.45
EfficientNetB0	F1-Score	0.0163	0.0087	0.0076	0.0250	2.12	7.33
EfficientNetB0	AUC-ROC	0.5523	0.0156	0.5367	0.5679	0.23	2.98

A.1.2 Paired Comparison Test Details

McNemar's Test Results

EfficientNetB0 vs ResNet50:

Chi-square statistic: 17.0

p-value: < 0.001

Effect size (Phi coefficient): 0.051

EfficientNetB0 vs MobileNetV2:

Chi-square statistic: 16.0

p-value: < 0.001

Effect size (Phi coefficient): 0.049

ResNet50 vs MobileNetV2:

Chi-square statistic: 5.12

p-value: 0.0234

Effect size (Phi coefficient): 0.028

A.1.3 Power Analysis Calculations

Table A.2: Statistical Power Analysis for Different Effect Sizes

Metric	Current Effect Size	Power ($\alpha=0.05$)	Min Detectable Effect	Required N for 80% Power
Sensitivity	0.0094	0.89	0.005	3,841
AUC-ROC	0.0468	0.84	0.025	4,112
F1-Score	0.0163	0.78	0.008	4,567

A.2 Computational Resource Analysis

A.2.1 Training Resource Detailed Breakdown

Table A.3: Comprehensive Training Resource Analysis

Architecture	GPU Hours	CPU Hours	Memory Peak (GB)	Energy (kWh)	CO2 Equivalent (kg)
EfficientNetB0	12.3	2.1	8.2	45.6	18.9
ResNet50	18.7	3.4	12.1	72.3	29.8
MobileNetV2	8.9	1.6	5.7	32.1	13.2

A.2.2 Inference Performance Metrics

Table A.4: Detailed Inference Performance Analysis

Architecture	GPU Inference (ms)	CPU Inference (ms)	Mobile CPU (ms)	Memory (MB)	FLOPS (Billion)
EfficientNetB0	23.4 ± 2.1	145.7 ± 12.3	892.3 ± 78.1	42.1	0.39
ResNet50	41.2 ± 3.8	267.9 ± 23.1	1,534.2 ± 134.7	97.3	4.09
MobileNetV2	12.8 ± 1.5	78.4 ± 8.9	456.7 ± 41.2	18.7	0.30

A.3 Clinical Performance Analysis

A.3.1 Confusion Matrix Details

Table A.5: Complete Confusion Matrix Analysis

Architecture	TP	FN	TN	FP	Sensitivity	Specificity	PPV	NPV
EfficientNetB0	1	105	6,503	16	0.0094	0.9975	0.0588	0.9840
ResNet50	0	106	6,495	24	0.0000	0.9963	0.0000	0.9840
MobileNetV2	0	106	6,512	7	0.0000	0.9989	0.0000	0.9840

A.3.2 ROC Curve Analysis

Table A.6: ROC Curve Coordinate Points

Architecture	Threshold	FPR	TPR	Precision	F1-Score
EfficientNetB0	0.1	0.0025	0.0094	0.0588	0.0163
EfficientNetB0	0.3	0.0015	0.0047	0.0476	0.0085
EfficientNetB0	0.5	0.0009	0.0000	0.0000	0.0000

A.4 Interpretability Analysis

A.4.1 Grad-CAM Attention Statistics

Table A.7: Quantitative Attention Pattern Analysis

Architecture	Mean Attention	Std Attention	Entropy	Gini Coefficient	Clinical Relevance Score
EfficientNetB0	0.0156	0.0234	4.23	0.76	0.14
ResNet50	0.0098	0.0145	5.89	0.85	0.08
MobileNetV2	0.0112	0.0198	5.34	0.82	0.09

A.4.2 Feature Attribution Analysis

Table A.8: ABCDE Criteria Correlation with AI Attention

Criteria	EfficientNetB0	ResNet50	MobileNetV2	Clinical Threshold
Asymmetry	0.23	0.12	0.15	0.70
Border	0.31	0.18	0.21	0.75
Color	0.28	0.14	0.19	0.65
Diameter	0.19	0.11	0.13	0.60
Evolution	N/A	N/A	N/A	0.80

Note: Correlation values range from 0 (no correlation) to 1 (perfect correlation).

Clinical thresholds represent minimum acceptable correlation for clinical relevance.

A.5 Implementation Guidelines

A.5.1 Recommended Evaluation Protocol

Standard Evaluation Checklist for Medical AI Systems:

1. Dataset Preparation

- Stratified sampling maintaining class distributions
- Multiple independent train/validation/test splits
- Demographic balance analysis and reporting
- Quality assessment and filtering procedures

2. Training Protocol Standardization

- Fixed random seeds for reproducibility
- Identical preprocessing across architectures
- Standardized hyper parameter optimization
- Multiple independent training runs

3. Performance Assessment

- Primary focus on clinically relevant metrics
- Comprehensive confidence interval reporting
- Statistical significance testing with appropriate corrections
- Effect size analysis and clinical significance assessment

4. Interpretability Evaluation

- Systematic attention pattern analysis
- Clinical expert validation of explanations
- Consistency assessment across cases
- Failure mode visualization and analysis

5. Computational Analysis

- Resource requirement documentation
- Scalability assessment
- Cost-effectiveness evaluation
- Deployment feasibility analysis

A.5.2 Clinical Deployment Readiness Assessment

Minimum Performance Thresholds for Clinical Consideration:

Metric	Screening Applications	Diagnostic Applications	Decision Support
Sensitivity	$\geq 85\%$	$\geq 90\%$	$\geq 80\%$
Specificity	$\geq 80\%$	$\geq 85\%$	$\geq 75\%$

Metric	Screening Applications	Diagnostic Applications	Decision Support
PPV	$\geq 25\%$	$\geq 50\%$	$\geq 20\%$
NPV	$\geq 99\%$	$\geq 99.5\%$	$\geq 98\%$
AUC-ROC	≥ 0.85	≥ 0.90	≥ 0.80

Additional Requirements:

- Clinical validation in prospective studies
- Multi-site performance verification
- Regulatory approval or clearance
- Comprehensive safety documentation

Provider training and certification programs