

DOCTORAL THESIS IN COMPUTER SCIENCE

Advancing Explainability and Predictive Accuracy in AI-Driven Cardiovascular Risk Assessment

A Hybrid Statistical and Neural Network Approach

Nishadha Himanshi Hikkaduwa Liyanage

FACULTY OF MATHEMATICS AND COMPUTER SCIENCE
UNIVERSITY OF ŁÓDŹ, POLAND

Submitted for the Degree of PhD in Computer Science

Supervisor: Prof. Andrzej Nowakowski

Co-supervisor: Dr Marta Lipnicka

Łódź, Poland

2026

Cardiovascular disease (CVD) remains the leading cause of mortality worldwide, creating a strong need for risk prediction models that are both accurate and interpretable in clinical settings. Traditional statistical methods are widely used in healthcare due to their transparency; however, they often struggle to capture the complex and high-dimensional nature of modern medical data. In contrast, machine learning approaches can achieve high predictive performance, but their limited interpretability restricts their practical use in clinical decision-making.

This thesis proposes a hybrid computational framework that combines data-driven modelling with interpretable risk estimation. The aim is to improve predictive performance while maintaining clarity and clinical relevance. By integrating feature selection, subgroup identification, and survival analysis, the proposed approach enables a more structured and personalised assessment of cardiovascular risk, supporting more informed clinical decision-making.

Abstract

Cardiovascular disease (CVD) remains a major global health challenge, making accurate and timely risk prediction essential for improving patient outcomes. Traditional statistical approaches, such as the Cox Proportional Hazards Model, provide interpretable results but are limited in their ability to capture complex and nonlinear relationships in clinical data. In contrast, machine learning models, particularly Feedforward Neural Networks (FFNNs), offer strong predictive performance but lack interpretability, which restricts their clinical applicability. This thesis investigates two complementary approaches for cardiovascular risk prediction. First, FFNN models are developed and evaluated to assess their ability to model nonlinear relationships in structured clinical data. Second, a novel hybrid computational framework is proposed, integrating Winner-Takes-All (WTA) competitive learning for patient clustering with cluster-specific Cox proportional hazards models. This approach enables subgroup-level risk estimation and captures heterogeneity within the patient population. To further enhance interpretability, a new cumulative risk measure, termed the Cumulative Prevalence Ratio (CPR), is introduced. CPR summarises time-dependent hazard information into a single interpretable value, facilitating direct comparison of long-term cardiovascular risk across patient clusters. The proposed methods are evaluated using a large-scale cardiovascular dataset, with performance assessed using the Concordance Index and ROC-based measures. Comparative analysis highlights the trade-off between the predictive strength of neural networks and the interpretability of the hybrid framework. The results indicate that the proposed hybrid approach achieves a balance between predictive performance and clinical interpretability. Overall, this work contributes to the development of transparent, subgroup-aware, and clinically meaningful models for cardiovascular risk prediction.

Keywords: Cardiovascular Disease, Risk Prediction, Survival Analysis, Neural Networks, Clustering, Interpretability, Hybrid Modelling

Acknowledgement

I would like to express my sincere gratitude to **Professor Andrzej Nowakowski**, my primary supervisor at the University of Łódź, whose insight, guidance, and feedback significantly shaped the direction and rigor of this research. I am also thankful to **Dr. Marta Lipnicka**, my co-supervisor, for her expertise and collaborative spirit, which enriched this work. I would also like to thank **Professor Stanisław Goldstein** for his support in helping initiate the process that led to the beginning of my Ph.D. studies. I am especially grateful to my parents for their unwavering love and encouragement, bridging the distance between Sri Lanka and Poland. Their sacrifices and belief in me provided the motivation to persevere and continue this journey. Finally, I would like to thank everyone—colleagues, friends, and faculty who supported me in different ways throughout my academic path. Your encouragement and kindness made this milestone possible.

Contents

1	INTRODUCTION	1
1.1	Background and Research Context	2
1.1.1	Research Gap	2
1.2	Problem Statement	3
1.2.1	Contributions of this Thesis	3
1.3	The Cardiovascular Disease (CVD) Dataset Used for the Research	4
1.3.1	Exploratory Data Analysis for Continuous Variables	5
1.3.2	Exploratory Data Analysis for Categorical Variables	10
1.4	Summary of Exploratory Data Analysis and Dataset Readiness	12
1.5	System Specification	13
1.5.1	Structure of the Thesis	13
2	APPLICATION OF FEEDFORWARD NEURAL NETWORKS FOR CARDIOVASCULAR RISK ASSESSMENT	15
2.1	Artificial Neural Networks (ANNs)	16
2.1.1	Theoretical Foundations of Artificial Neural Networks	16
2.1.2	Applications and Types of Artificial Neural Networks	17
2.1.3	The Perceptron	17
2.1.4	Perceptron Learning Algorithm	19
2.1.5	Example Calculation with Perceptron Learning	20
2.1.6	Perceptron Algorithm Training Process	22
2.2	Feedforward Neural Networks for Cardiovascular Risk Prediction	23
2.2.1	Activation Function and Output Node	24
2.2.2	Weights and Biases	24
2.2.3	Model Architecture and Training Setup	25
2.2.4	Model Selection and Architecture Visualisation	27
2.2.5	Evaluation Metrics	27
2.2.6	Architecture Performance Analysis	28
2.2.7	Accuracy Trend Across Architectures	29
2.2.8	Best Model Performance	30
2.2.9	Prediction Process and Model Output	30
2.2.10	Results and Performance Analysis	31
2.3	Discussion	34
2.4	Conclusion	35
3	A HYBRID COMPUTATIONAL FRAMEWORK FOR CARDIOVASCULAR RISK STRATIFICATION	36
3.1	Introduction	37
3.2	Related Work	39
3.3	Problem Formulation and Research Objectives	41
3.4	Proposed Computational Framework	43
3.4.1	Stage 1: WTA-Based Clustering	43

3.4.2	Stage 2: Cluster-Specific Survival Modelling	45
3.4.3	Stage 3: Cumulative Risk Estimation using CPR	45
3.4.4	Risk Stratification	46
3.4.5	Experimental Setup	47
3.5	Results	48
3.5.1	WTA-Based Clustering Results	49
3.5.2	Cluster-Specific Survival Analysis	53
3.5.3	Prediction for New Patients	55
3.5.4	Prediction Results	57
3.5.5	Extended Evaluation of Calibration, Risk Stratification, and Structural Performance	59
3.5.6	Summary of Results	62
3.6	Discussion	63
3.6.1	Strengths, Limitations, and Future Directions	65
4	DISCUSSION AND CONCLUSION	67
4.1	Discussion	68
4.1.1	Addressing the Research Problem	68
4.1.2	Comparison Between the FFNN Baseline and the Hybrid Framework	68
4.1.3	ROC Curve Comparison	71
4.1.4	Interpretability Comparison	72
4.1.5	Importance of Modelling Patient Heterogeneity	73
4.1.6	Temporal Risk Representation and CPR	73
4.1.7	Clinical and Practical Implications	74
4.1.8	Study-Level Evidence Positioning Using Published Cardiovascular Prediction Studies	75
4.1.9	Summary of Discussion	78
4.2	Conclusion	78
4.2.1	Future Directions	79
	References	i
	Appendix	vii

List of Figures

1.1	Histogram and KDE plot for age distribution	6
1.2	Boxplot for age distribution	7
1.3	Histogram and KDE plot for body mass index distribution	7
1.4	Boxplot for body mass index distribution	8
1.5	Histogram and KDE plot for diastolic blood pressure	8
1.6	Boxplot for diastolic blood pressure	9
1.7	Histogram and KDE plot for fasting blood sugar	9
1.8	Boxplot for fasting blood sugar	10
1.9	Distribution of categorical variables	10
1.10	CVD prevalence across categorical variables	11
1.11	Cramér’s V heatmap for categorical variables	12
2.1	Single-layer perceptron model with two inputs and one output neuron	18
2.2	Geometric representation of the perceptron decision boundary in two-dimensional space.	19
2.3	Best-performing FFNN architecture selected based on experimental evaluation. The network consists of nine input features, three hidden layers with ReLU activation, and a single output neuron with sigmoid activation for CVD probability estimation.	27
2.4	Accuracy comparison across FFNN architectures	30
2.5	Confusion matrix for the best FFNN model	32
2.6	Training and validation loss curves	33
2.7	ROC curve for the best FFNN model	34
3.1	Silhouette score across different numbers of WTA clusters.	49
3.2	PCA-based visualisation of WTA clusters.	51
3.3	t-SNE visualisation of WTA clusters.	52
3.4	Kaplan–Meier survival curves for selected clusters (Kaplan & Meier, 1958). . .	54
3.5	Prediction process for new patients. A new patient is first assigned to a subgroup using the WTA assignment rule, followed by cluster-specific CPR estimation and final cardiovascular risk stratification.	56
3.6	Distribution of predicted risk categories	57
3.7	Calibration plot of the hybrid cardiovascular risk framework. The curve shows strong agreement between predicted and observed event probabilities, with only minor deviation in lower-risk regions.	60
3.8	Observed cardiovascular event rates across predicted risk categories. A clear upward trend is visible, with a sharp increase in the high-risk group.	61

4.1	ROC curve comparison between the FFNN model and the proposed hybrid framework.	71
4.2	Forest-style evidence plot of cardiovascular prediction AUC values.	75

List of Tables

1.1	Descriptive statistics of continuous variables in the CVD SL dataset	5
1.2	Chi-square test results for categorical variables	12
2.1	Sample data for perceptron training	20
2.2	Performance comparison of FFNN architectures (best model highlighted) . . .	29
3.1	Cluster size distribution after WTA clustering	50
3.2	Clustering validation metrics	51
3.3	Cluster-wise Cox model fitting status	53
3.4	Distribution of predicted risk categories	57
3.5	Cluster-level prediction summary	58
3.6	Representative prediction examples	59
3.7	Calibration performance of the hybrid cardiovascular risk framework.	59
3.8	Observed cardiovascular event rates across predicted risk categories.	61
4.1	Performance comparison between the FFNN baseline model and the proposed hybrid framework evaluated on unseen cardiovascular test data.	69
4.2	Interpretability comparison between the FFNN model and the proposed hybrid framework	72
4.3	Contextual comparison of published cardiovascular prediction studies used for evidence positioning.	77

CHAPTER 1

INTRODUCTION

1.1 Background and Research Context

Cardiovascular disease (CVD) remains one of the leading causes of morbidity and mortality worldwide, accounting for approximately 17.9 million deaths each year (World Health Organization, 2021). It includes a range of conditions, such as coronary artery disease, stroke, heart failure, and peripheral vascular disease. Despite continuous advances in prevention and treatment, the global burden of CVD remains high, particularly in ageing populations and resource-limited healthcare systems (Benjamin et al., 2017). Accurate risk prediction plays a critical role in early intervention and personalised treatment planning. Traditional statistical approaches, including logistic regression and the Cox Proportional Hazards Model (CPHM), are widely used in clinical practice due to their interpretability and strong theoretical foundation (Harrell, 2015). Clinical tools such as the Framingham Risk Score are based on well-established risk factors, including age, blood pressure, cholesterol levels, smoking status, and diabetes (Wilson et al., 1998). These models provide clear and understandable outputs, which support their adoption in healthcare settings.

However, traditional models rely on simplifying assumptions, such as linear relationships between variables and population homogeneity. These assumptions may limit their ability to fully capture the complexity and variability of real-world clinical data (DeFilippis et al., 2015). As a result, there is increasing interest in more flexible modelling approaches that can better represent nonlinear interactions and high-dimensional feature spaces. Recent developments in machine learning and deep learning have provided new opportunities for improving predictive performance. In particular, Feedforward Neural Networks (FFNNs) (see Chapter 2) have demonstrated strong capability in modelling complex relationships within structured clinical datasets (Krittanawong et al., 2020). Despite these advantages, their lack of transparency remains a major limitation, as it restricts their practical use in clinical decision-making where interpretability and trust are essential (Ching et al., 2018; Rudin, 2019). To address this challenge, hybrid modelling approaches have emerged, aiming to combine the strengths of data-driven learning with the interpretability of traditional statistical methods. These approaches seek to improve predictive performance while maintaining clinical relevance and transparency. Building on this direction, this thesis proposes a hybrid computational framework (see Chapter 3) that integrates clustering with survival analysis. By identifying patient subgroups and modelling risk within more homogeneous contexts, the framework enables a more structured and context-aware representation of cardiovascular risk. In addition, a cumulative risk measure is introduced to support the assessment of long-term outcomes. This integrated approach aims to provide both accurate and interpretable risk predictions, forming the basis for the methods developed in the following chapters.

1.1.1 Research Gap

Despite significant progress in cardiovascular risk prediction, several important limitations remain. Traditional statistical models are constrained in their ability to capture complex

and nonlinear clinical relationships, while machine learning approaches, although often more accurate, tend to lack interpretability, which limits their adoption in clinical practice (Doshi-Velez & Kim, 2017; Harrell, 2015; Rudin, 2019). Existing hybrid approaches attempt to address this trade-off; however, many still rely on global modelling strategies that do not explicitly account for subgroup-specific variation. As a result, they may fail to capture differences in risk across heterogeneous patient populations. Furthermore, current methods often lack a clear and interpretable representation of long-term risk, making it difficult to compare outcomes across different patient groups. These limitations highlight the need for a unified modelling framework that can capture complex relationships in clinical data, account for variability across patient subgroups, provide interpretable risk estimates, and support a meaningful representation of cumulative cardiovascular risk. Addressing these requirements forms the foundation of the research problem considered in this thesis.

1.2 Problem Statement

Current approaches to cardiovascular risk prediction do not simultaneously achieve strong predictive performance and clinically interpretable outputs. This limitation reduces their practical usefulness in real-world healthcare settings, where both accuracy and transparency are essential. To address this challenge, this thesis develops a hybrid modelling framework that integrates clustering and survival analysis to produce structured and interpretable risk estimates. The study is guided by the following research questions:

- How does the predictive performance of FFNNs (see Chapter 2) compare with the proposed hybrid framework (see Chapter 3)?
- Can subgroup-based modelling improve interpretability without reducing predictive accuracy?
- What trade-offs arise between global FFNN models (Chapter 2) and the hybrid approach (Chapter 3) in clinical prediction tasks?
- How can model outputs from both approaches be structured to better support clinical decision-making (see Chapters 2 and 3)?

These questions guide the design, evaluation, and comparison of the modelling approaches presented in this thesis.

1.2.1 Contributions of this Thesis

This thesis makes several key contributions to the field of cardiovascular risk prediction. First, it proposes a hybrid computational framework (see Chapter 3) that integrates Winner-Takes-All (WTA) clustering with cluster-specific Cox proportional hazards models. This approach captures subgroup-level variation in risk and provides interpretable predictions aligned with clinical reasoning. Second, the thesis introduces the Cumulative Prevalence Ratio (CPR), a novel metric that summarises time-dependent risk into a single interpretable

value. This enables consistent comparison of long-term cardiovascular risk across patient subgroups. Third, a comparative evaluation is conducted between the FFNN baseline model (see Chapter 2) and the proposed hybrid framework using metrics such as the Concordance Index (C-index) and AUC–ROC. This analysis highlights the trade-off between predictive accuracy and interpretability. Finally, this work contributes to the development of clinically applicable and interpretable artificial intelligence models by combining data-driven learning with structured risk modelling. The proposed framework enhances transparency, supports decision-making, and improves the practical usability of predictive models in healthcare settings.

1.3 The Cardiovascular Disease (CVD) Dataset Used for the Research

This study uses a large-scale dataset, referred to as *CVD SL*, which contains medical records collected from three healthcare institutions in Sri Lanka, including both government hospitals and private nursing homes. The dataset consists of 66,816 patient observations, each described by a set of demographic, clinical, and lifestyle variables. Before analysis, the dataset was systematically examined for missing values, inconsistencies, and data quality issues. Records with incomplete or invalid entries were removed to ensure the reliability of the modelling process. In addition to the binary CVD outcome, time-to-event information was derived from patient records, enabling the application of survival analysis methods in later stages of this study.

The dataset includes the following key variables:

- **Age:** A numerical variable representing the patient’s age in years.
- **Body Mass Index (BMI):** A continuous variable representing body mass index.
- **Gender:** A categorical variable originally coded as 1 for males and 2 for females. For modelling purposes, this variable was recoded into binary format, where 1 represents males and 0 represents females.
- **Diastolic Blood Pressure (DBP):** A clinical variable representing diastolic blood pressure.
- **Cholesterol Level:** A categorical variable defined as:
 - 1 – Normal
 - 2 – Above normal
 - 3 – Well above normal
- **Smoking Habit:** A binary variable where 1 indicates a smoker and 0 indicates a non-smoker.
- **Alcohol Intake:** A binary variable where 1 indicates alcohol consumption and 0 indicates no alcohol consumption.

- **Physical Activity:** A binary variable where 1 indicates an active individual and 0 indicates an inactive individual.
- **Fasting Blood Sugar (FBS):** A numerical variable representing fasting blood sugar levels.
- **Presence of Cardiovascular Disease (CVD):** The target variable:
 - 1 – Presence of CVD
 - 0 – Absence of CVD

These variables provide a comprehensive representation of patient health status and cardiovascular risk. The combination of demographic, clinical, and lifestyle factors supports detailed analysis of relationships associated with cardiovascular disease. The size and diversity of the dataset make it suitable for both statistical modelling and machine learning applications. With a large number of observations and heterogeneous features, the dataset enables the identification of complex patterns and variation across different patient groups. This strengthens the robustness and generalisability of the proposed models and provides a solid foundation for developing reliable and clinically relevant cardiovascular risk prediction systems.

1.3.1 Exploratory Data Analysis for Continuous Variables

Table 1.1: Descriptive statistics of continuous variables in the CVD SL dataset

Variable	Mean	Std Dev	Min – Max	Median (IQR)
Age (Years)	53.14	7.14	39 – 65	54 (48, 59)
BMI	32.12	9.55	16 – 45	32.47 (23.73, 41.19)
Diastolic Blood Pressure (mmHg)	89.32	9.53	70 – 110	85 (85, 95)
Fasting Blood Sugar (mg/dL)	127.36	16.27	86 – 161	128 (120, 137)

Exploratory data analysis (EDA) is a fundamental step in statistical and machine learning workflows, enabling the identification of key patterns, anomalies, and distributional characteristics within a dataset (Tukey, 1977). In medical research, EDA supports a deeper understanding of data structure and improves the reliability and interpretability of predictive modelling approaches (Komorowski et al., 2018). The integration of EDA with machine learning has been widely recognised as essential in clinical and epidemiological studies (James et al., 2021). In this study, EDA was conducted on the CVD SL dataset to examine the distribution of continuous variables, assess variability, and identify potential outliers. Visualisation techniques, including histograms, kernel density estimation (KDE) plots, boxplots, and correlation analysis, were systematically applied to explore data patterns and relationships. Table 1.1 presents key descriptive statistics for four continuous variables:

age, body mass index (BMI), diastolic blood pressure (DBP), and fasting blood sugar (FBS). The dataset contains 66,816 patient records, providing a sufficiently large sample for reliable statistical analysis. These descriptive statistics provide an initial overview of the dataset. The age distribution reflects a predominantly middle-aged population with moderate variability. BMI shows substantial variation, indicating the presence of different weight categories within the dataset. DBP values suggest variability in cardiovascular health status, while FBS levels indicate differences in glycaemic control across individuals. These observations provide a foundation for further analysis and inform feature selection and model development in later stages of this study.

Interpretation of Exploratory Data Analysis

This section interprets the continuous variables based on their distributions, variability, and the presence of outliers. Understanding these characteristics is essential for identifying underlying data patterns and ensuring the reliability of subsequent predictive modelling.

Age Distribution

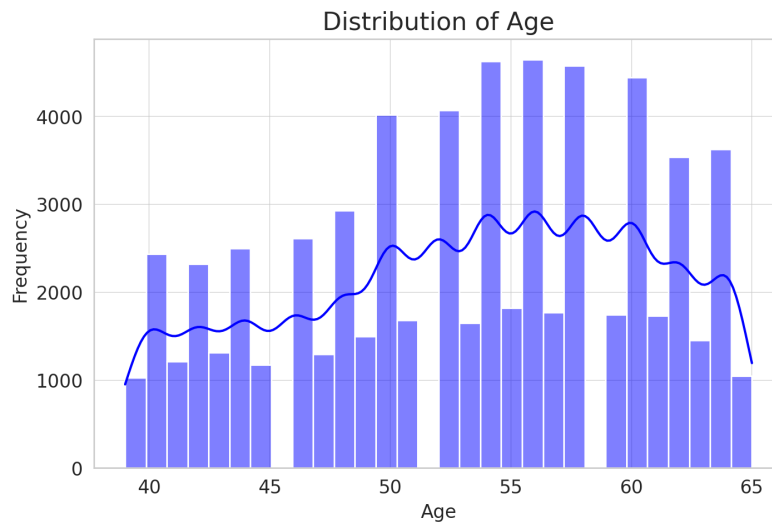


Figure 1.1: Histogram and KDE plot for age distribution

The histogram and KDE plot in Figure 1.1 show a unimodal distribution, with most observations concentrated between 40 and 65 years. The distribution exhibits a slight right skew, indicating a higher proportion of individuals in older age groups. Such patterns are commonly observed in epidemiological datasets, where ageing populations are associated with increased risk of chronic diseases (World Health Organization, 2021). The boxplot in Figure 1.2 further illustrates the spread of the data, with an interquartile range approximately between 48 and 60 years and a median near 54 years. The absence of extreme outliers suggests that the sample is stable and representative, reducing the risk of bias in subsequent modelling (Tukey, 1977).

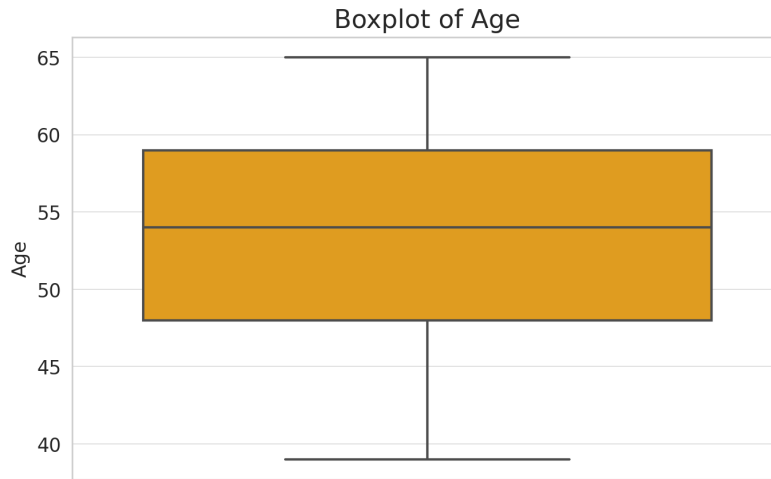


Figure 1.2: Boxplot for age distribution

Body Mass Index Distribution

The histogram in Figure 1.3 shows a clear right-skewed distribution, with a primary concentration of values around 30 and a secondary peak near 45. This indicates a relatively high prevalence of elevated BMI values within the population. Such right-skewed patterns are typical in population health data, where overweight and obesity are increasingly common (Ng et al., 2014). The secondary peak may be influenced by rounding or recording effects, which are frequently observed in clinical datasets.

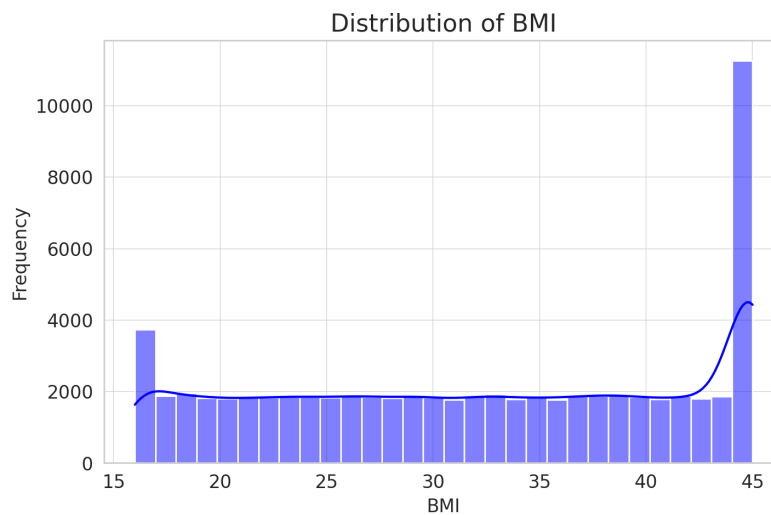


Figure 1.3: Histogram and KDE plot for body mass index distribution

The boxplot in Figure 1.4 highlights substantial variability, with an interquartile range approximately between 25 and 40 and a median around 32. Several high-value outliers are present, indicating extreme BMI values. These may correspond to individuals with severe obesity, which is a well-established risk factor for cardiovascular disease and may require careful handling during modelling (Hruby & Hu, 2015)

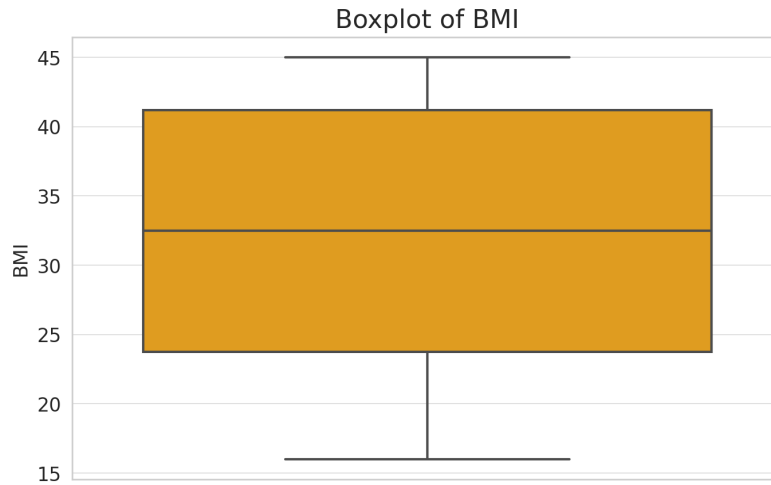


Figure 1.4: Boxplot for body mass index distribution

Diastolic Blood Pressure Distribution

The histogram and KDE curve in Figure 1.5 show a multimodal pattern, with peaks around 75, 85, and 95 mmHg. This likely reflects standardised clinical measurement ranges and the presence of distinct subgroups within the dataset. Multimodal distributions often indicate underlying population heterogeneity or discretisation effects in recorded measurements (Silverman, 1986).

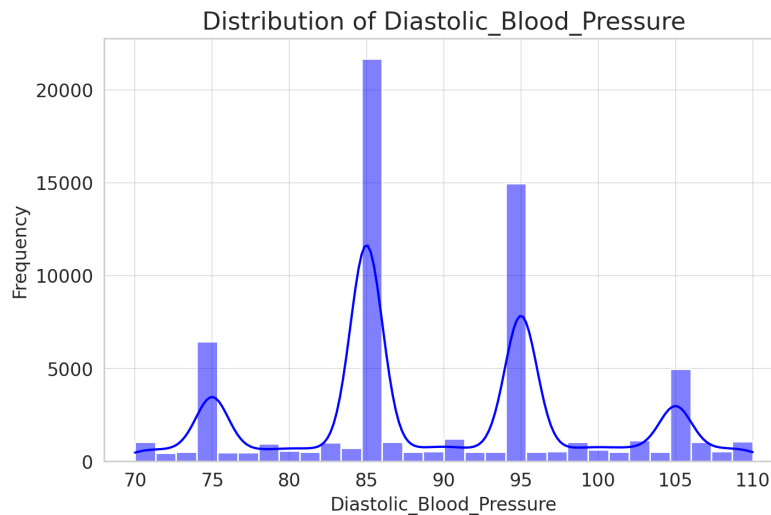


Figure 1.5: Histogram and KDE plot for diastolic blood pressure

The boxplot in Figure 1.6 supports this observation, with an interquartile range between 80 and 95 mmHg and a median near 90 mmHg. No extreme outliers are observed, and all values remain within clinically plausible ranges. This suggests consistent measurement quality and reduces the likelihood of data entry errors.

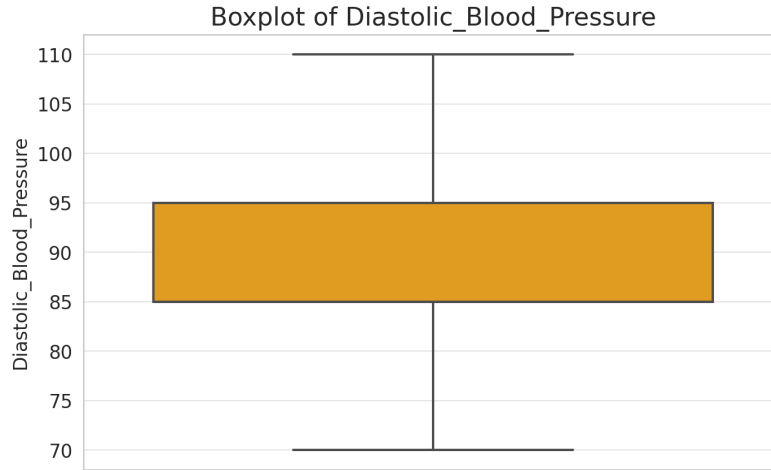


Figure 1.6: Boxplot for diastolic blood pressure

Fasting Blood Sugar Distribution

The histogram in Figure 1.7 indicates a right-skewed distribution, with most values ranging between 100 and 140 mg/dL. The KDE curve peaks around 130 mg/dL, suggesting that a substantial proportion of individuals fall within prediabetic or diabetic ranges. This pattern aligns with known trends in population-level glycaemic measurements, where elevated glucose levels are increasingly prevalent (American Diabetes Association, 2022).

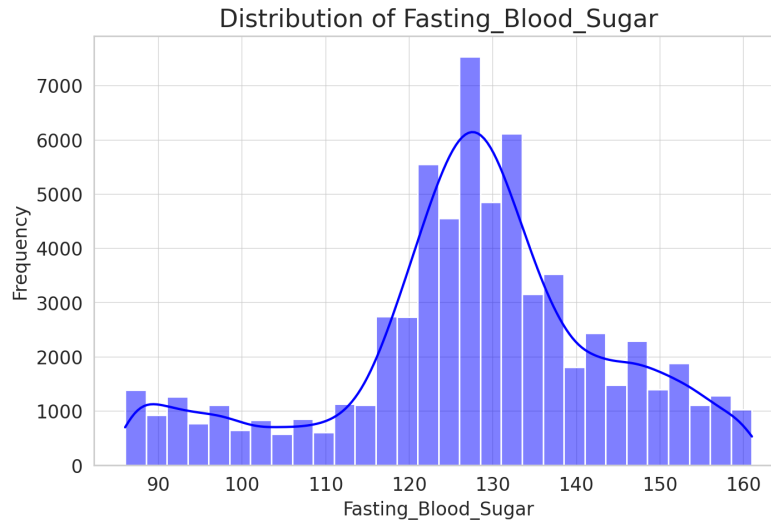


Figure 1.7: Histogram and KDE plot for fasting blood sugar

The boxplot in Figure 1.8 further confirms this pattern, with an interquartile range approximately between 110 and 140 mg/dL and a median near 125 mg/dL. Several high-value outliers are present, indicating significantly elevated glucose levels. These extreme values may correspond to undiagnosed or poorly controlled diabetes and are important for risk modelling and stratification.

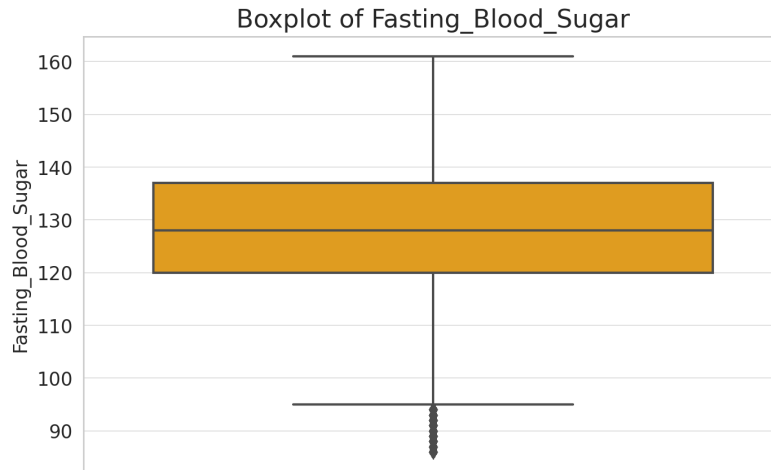


Figure 1.8: Boxplot for fasting blood sugar

1.3.2 Exploratory Data Analysis for Categorical Variables

In addition to continuous variables (see Section 1.3.1), categorical variables were analysed to examine their distribution and relationship with cardiovascular disease (CVD). The variables considered include gender, smoking habit, alcohol intake, physical activity, and cholesterol level. Bar charts were used to visualise overall distributions, while stacked bar charts were applied to assess their association with CVD prevalence (Tukey, 1977).

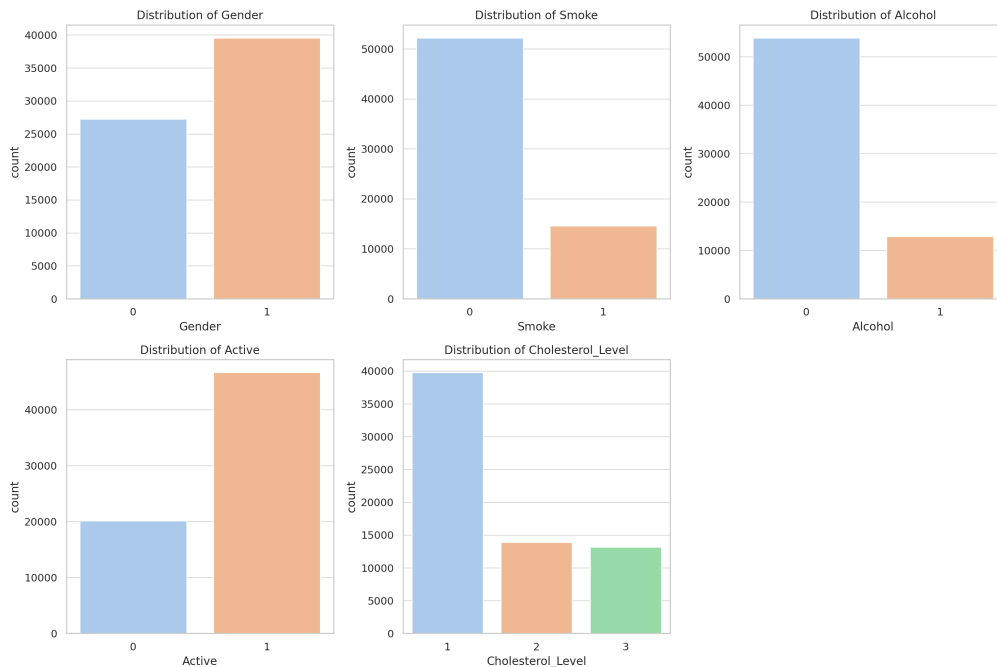


Figure 1.9: Distribution of categorical variables

Figure 1.9 shows the distribution of categorical variables across the dataset. A higher proportion of males is observed compared to females, indicating a class imbalance that

may influence risk patterns. Most individuals are non-smokers, while alcohol consumption is reported in a smaller subset of the population. A larger proportion of individuals are physically active, and most patients fall into the normal cholesterol category, with fewer cases in elevated groups. These patterns are consistent with typical epidemiological datasets, where behavioural and lifestyle variables are unevenly distributed (James et al., 2021). To further examine relationships with CVD, stacked bar charts were generated, as shown in Figure 1.10. Figure 1.10 indicates that CVD prevalence varies across categories. A higher proportion of CVD cases is observed among males, suggesting a potential influence of gender, which is consistent with established cardiovascular risk patterns (Benjamin et al., 2017). Smoking is associated with increased CVD prevalence, supporting its well-established role as a major cardiovascular risk factor (World Health Organization, 2021). Alcohol consumption shows a slight increase in CVD cases, although the difference is less pronounced. In contrast, physically active individuals tend to exhibit lower CVD prevalence, reflecting the protective effect of physical activity (Warburton et al., 2006). Cholesterol level shows a clear gradient, with higher CVD prevalence observed in elevated categories, aligning with clinical evidence linking dyslipidaemia to cardiovascular risk (Ference et al., 2017).

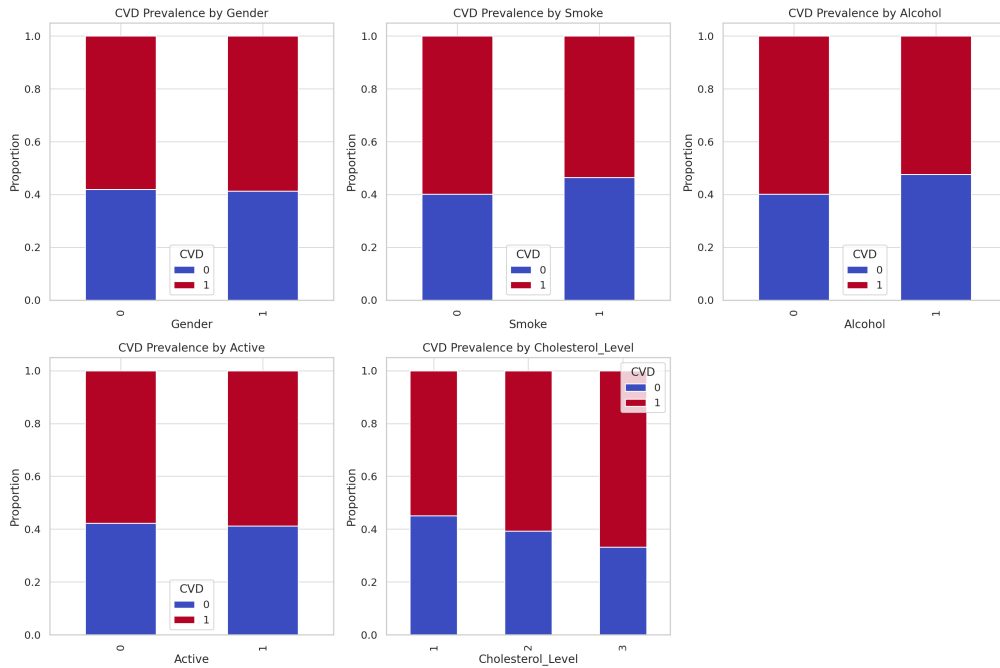


Figure 1.10: CVD prevalence across categorical variables

To evaluate statistical significance, a chi-square test was conducted, and the results are summarised in Table 1.2. All variables show statistically significant associations with CVD. However, statistical significance does not indicate the strength of these relationships. To address this, Cramér's V was used to quantify association strength. As shown in Figure 1.11, all variables exhibit weak associations with CVD, with values close to zero. Physical activity shows the strongest relationship, although the effect remains small. This is consistent with standard interpretations of effect size, where small values indicate limited

practical impact despite statistical significance (Akoglu, 2018; Cohen, 1988).

Table 1.2: Chi-square test results for categorical variables

Categorical Variable	Chi-square p-value
Gender	0.0001
Smoking	< 0.0001
Alcohol Intake	0.0023
Physical Activity	< 0.0001
Cholesterol Level	< 0.0001

Overall, these findings suggest that while categorical variables are statistically associated with CVD, their individual predictive strength is limited. This highlights the importance of modelling interactions between variables and identifying subgroup-level patterns, which are explored further in later stages of this study.

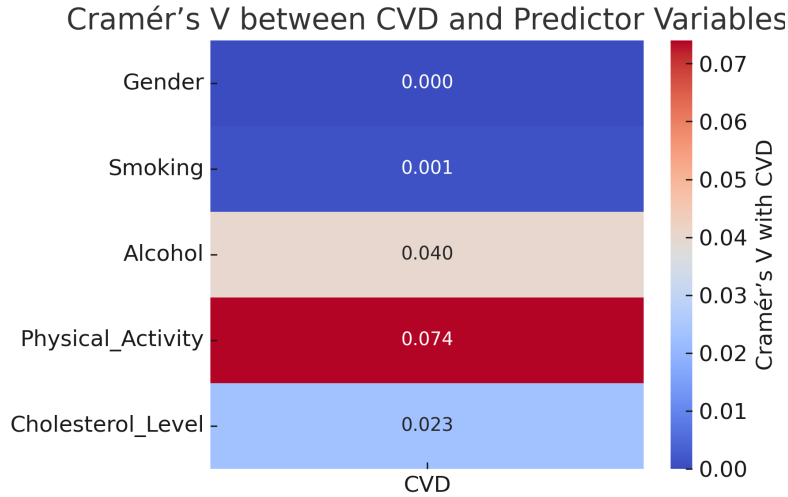


Figure 1.11: Cramér's V heatmap for categorical variables

1.4 Summary of Exploratory Data Analysis and Dataset Readiness

This study utilises the CVD SL dataset, which consists of 66,816 patient records collected from healthcare institutions in Sri Lanka. The dataset includes demographic, clinical, and lifestyle variables relevant to cardiovascular disease (CVD) risk. Exploratory data analysis (EDA) was conducted to examine data distribution, variability, and overall quality, providing a foundation for subsequent predictive modelling (James et al., 2021; Tukey, 1977). The analysis of continuous variables (see Section 1.3.1) revealed several important patterns. Age follows a slightly right-skewed distribution, indicating a predominantly middle-aged population. Body mass index (BMI) shows high variability with a shift towards elevated values, suggesting a notable presence of overweight and obese individuals.

Diastolic blood pressure (DBP) exhibits a multimodal distribution, which is common in clinical measurements, while fasting blood sugar (FBS) shows right skewness, indicating that a subset of patients may have impaired glycaemic control (Hastie et al., 2009). Categorical variables (see Section 1.3.2), including gender, smoking habit, alcohol intake, physical activity, and cholesterol level, were analysed to assess their association with CVD. Chi-square tests confirmed statistically significant relationships ($p < 0.05$) between these variables and disease prevalence. However, while these relationships are statistically significant, Cramér’s V analysis indicates that their effect sizes are weak across all variables (Akoglu, 2018). Physical activity shows the strongest association, although the magnitude remains small (Cramér’s V ≈ 0.074). Other variables, such as smoking and gender, show minimal practical influence despite statistical significance. This distinction between statistical significance and effect size is important in predictive modelling. In large datasets, even small effects may appear statistically significant without providing meaningful predictive value (Cohen, 1988). The observed weak associations suggest that individual variables alone are insufficient for accurate risk prediction. Instead, modelling approaches that capture interactions between variables and account for patient heterogeneity are required. Overall, the dataset demonstrates high completeness and sufficient variability, indicating its suitability for further modelling. The EDA findings provide insight into the structure and relationships within the data, supporting feature selection and guiding the development of models that aim to balance predictive performance and interpretability.

1.5 System Specification

All computational experiments and data analyses in this study were conducted on a standard personal computing system. The hardware configuration consisted of Windows 11 Home (64-bit, Version 25H2), powered by an Intel Core Ultra 5 226V processor (2.10 GHz), with 16 GB RAM and 954 GB SSD storage. The software environment was based on Python 3.13.5, implemented using Jupyter Notebook as the development platform. Key libraries included NumPy 2.1.3, Pandas 2.2.3, Scikit-learn 1.6.1, Statsmodels 0.14.4, and Matplotlib 3.10.0, which were used for data preprocessing, statistical analysis, model development, and visualisation. This environment provided sufficient computational capacity to efficiently perform all stages of the proposed methodology, including data preprocessing, exploratory data analysis, model training, and evaluation. Although the experiments were conducted on a standard system, the results demonstrate that the proposed approach is computationally efficient and does not require specialised high-performance hardware. This supports its applicability in real-world clinical and research settings.

1.5.1 Structure of the Thesis

The remainder of this thesis is organised as follows:

- **Chapter 2: Application of Feedforward Neural Networks for Cardiovascular Risk Prediction**

This chapter presents a baseline modelling approach based on Feedforward Neural Networks (FFNNs). It describes the model architecture, training process, and evaluation using standard classification metrics. The results provide a reference point for assessing predictive performance and highlight the characteristics of global models applied to clinical data.

- **Chapter 3: A Hybrid Computational Framework for Cardiovascular Risk Assessment**

This chapter introduces the proposed hybrid framework, integrating clustering, survival modelling, and cumulative risk estimation within a unified approach. Patient subgroups are identified using a Winner-Takes-All (WTA) mechanism, followed by cluster-specific Cox proportional hazards models. Long-term risk is summarised using the Cumulative Prevalence Ratio (CPR), enabling structured analysis across patient groups and over time.

- **Chapter 4: Discussion and Conclusion**

This chapter synthesises the findings from the FFNN baseline and the proposed hybrid framework. It provides a comparative analysis of their performance and modelling characteristics, focusing on how differences in model structure influence predictive behaviour. The chapter also evaluates the strengths and limitations of the approach, relates the results to existing research, and outlines directions for future work.

Overall, this thesis aims to bridge the gap between predictive performance and clinical usability by developing models that are both accurate and interpretable for cardiovascular risk assessment.

CHAPTER 2

APPLICATION OF FEEDFORWARD NEURAL NETWORKS FOR CARDIOVASCULAR RISK ASSESSMENT

2.1 Artificial Neural Networks (ANNs)

Artificial Neural Networks (ANNs) are computational models inspired by the structure and function of biological neural systems. They are designed to learn complex patterns directly from data through iterative training processes. Owing to their flexibility and strong predictive capabilities, ANNs have become widely used in modern data analysis, particularly in medical diagnostics and clinical decision support systems (Esteva et al., 2019; Topol, 2019). In the context of cardiovascular disease (CVD) risk prediction, ANNs are particularly useful because they can model nonlinear relationships between multiple clinical variables, which are often difficult to capture using traditional statistical approaches.

2.1.1 Theoretical Foundations of Artificial Neural Networks

The development of Artificial Neural Networks (ANNs) can be traced back to the computational neuron model proposed by McCulloch and Pitts in 1943, which aimed to represent fundamental aspects of neural activity using mathematical logic (McCulloch & Pitts, 1943). This early work laid the foundation for modern neural network architectures composed of interconnected processing units capable of learning from data (Goodfellow et al., 2016).

In its simplest form, an artificial neuron computes a weighted sum of input features followed by a nonlinear transformation. Let the input vector be defined as $\mathbf{x} \in \mathbb{R}^n$, the weight vector as $\mathbf{w} \in \mathbb{R}^n$, and the bias term as $b \in \mathbb{R}$. The neuron computes the weighted input as

$$z = \mathbf{w}^T \mathbf{x} + b$$

where $\mathbf{w}^T \mathbf{x} \in \mathbb{R}$ denotes the inner product between the transpose of the weight vector and the input vector, and $z \in \mathbb{R}$ represents the linear combination of the input features. An activation function $\sigma(\cdot)$ is then applied to transform the weighted input into the neuron output:

$$a = \sigma(z)$$

where $a \in \mathbb{R}$ represents the output activation of the neuron and $\sigma(\cdot)$ denotes a nonlinear activation function. During training, the network iteratively updates its parameters, including weights and biases, to minimise a predefined loss function. This optimisation process enables the model to generalise to previously unseen data. Such capability is particularly important in medical applications, where accurate prediction and classification are essential for clinical decision-making. However, despite their strong predictive performance, neural network models do not inherently provide transparency in how predictions are generated. This limitation reduces interpretability in clinical settings, where understanding the reasoning behind predictions is essential for trust and practical adoption.

2.1.2 Applications and Types of Artificial Neural Networks

Artificial Neural Networks (ANNs) have been successfully applied across various domains, including healthcare, engineering, and finance. In cardiovascular medicine, they have demonstrated strong performance in risk prediction and diagnostic tasks (Krittawong et al., 2020). Training an ANN involves adjusting weights and biases using optimisation algorithms to minimise a predefined loss function. In general, larger and more diverse datasets improve model performance by enabling more robust parameter estimation (Goodfellow et al., 2016). Several ANN architectures have been developed to address different types of data and tasks:

- **Feedforward Neural Networks (FFNNs):** Also known as multilayer perceptrons, these represent the simplest class of neural networks and are widely used for classification and regression tasks, particularly with structured tabular data such as clinical datasets (Rumelhart et al., 1986).
- **Convolutional Neural Networks (CNNs):** Designed for grid-structured data such as images (LeCun et al., 2015).
- **Recurrent Neural Networks (RNNs):** Designed for sequential and temporal data (Hochreiter & Schmidhuber, 1997).

Among these architectures, FFNNs are the most appropriate for this study, as the dataset consists of structured clinical variables rather than image or sequential data.

2.1.3 The Perceptron

The perceptron is one of the earliest supervised learning models and is primarily used for binary classification tasks (Rosenblatt, 1958). It computes a linear combination of input features and applies a threshold activation function to produce a binary output. Given an input vector

$$\mathbf{x} = [x_1, x_2, \dots, x_n]^T \in \mathbb{R}^n$$

and a corresponding weight vector

$$\mathbf{w} = [w_1, w_2, \dots, w_n]^T \in \mathbb{R}^n,$$

the perceptron output $y \in \{0, 1\}$ is defined as:

$$y = F\left(\sum_{i=1}^n w_i x_i + \theta\right) \quad (2.1)$$

where:

- $x_i \in \mathbb{R}$ denotes the i -th input feature,
- $w_i \in \mathbb{R}$ denotes the weight associated with input feature x_i ,

- $\theta \in \mathbb{R}$ is the bias term,
- $n \in \mathbb{N}$ represents the number of input features,
- $y \in \{0, 1\}$ represents the predicted binary class label.

The quantity

$$s = \sum_{i=1}^n w_i x_i + \theta$$

represents the weighted input (linear combination) to the activation function, where $s \in \mathbb{R}$.

The activation function $F(s)$ is defined as:

$$F(s) = \begin{cases} 1, & \text{if } s \geq 0 \\ 0, & \text{if } s < 0 \end{cases} \quad (2.2)$$

This formulation defines a linear decision boundary, allowing the perceptron to separate linearly separable data.

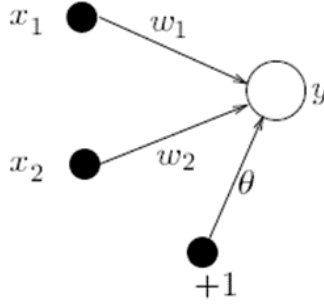


Figure 2.1: Single-layer perceptron model with two inputs and one output neuron

For a simple case with two input features, the decision boundary is defined as:

$$w_1 x_1 + w_2 x_2 + \theta = 0 \quad (2.3)$$

which can be rewritten as:

$$x_2 = -\frac{w_1}{w_2} x_1 - \frac{\theta}{w_2} \quad (2.4)$$

where:

- $x_1, x_2 \in \mathbb{R}$ are the input features,
- $w_1, w_2 \in \mathbb{R}$ are the corresponding feature weights,
- $\theta \in \mathbb{R}$ is the bias term.

Equation 2.4 represents a straight-line decision boundary in two-dimensional space, separating the input space into two regions corresponding to the two output classes. These formulations provide the conceptual basis for understanding more complex neural network architectures,

where multiple perceptrons are combined to model nonlinear relationships and hierarchical feature representations (Goodfellow et al., 2016; Minsky & Papert, 1969).

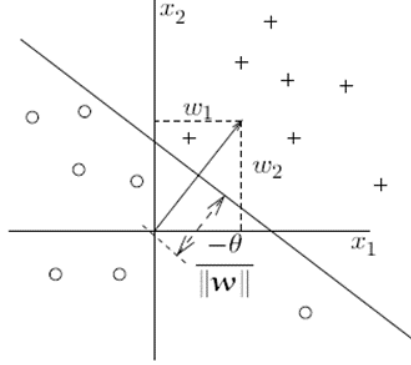


Figure 2.2: Geometric representation of the perceptron decision boundary in two-dimensional space.

2.1.4 Perceptron Learning Algorithm

Building on the perceptron model described above, the learning process determines how the model parameters are updated during training. The perceptron learning algorithm is a supervised learning method that updates weights based on prediction error. For each training sample, the model computes a predicted output $\hat{y} \in \{0, 1\}$ using Equation 2.1 and compares it with the true target value.

The prediction error is defined as:

$$e = y_{\text{true}} - \hat{y} \quad (2.5)$$

where:

- $e \in \mathbb{R}$ denotes the prediction error,
- $y_{\text{true}} \in \{0, 1\}$ denotes the true class label,
- $\hat{y} \in \{0, 1\}$ denotes the predicted class label produced by the perceptron.

Model parameters are updated using the perceptron learning rule (Haykin, 1999):

$$w_i(t+1) = w_i(t) + \eta e x_i \quad (2.6)$$

$$\theta(t+1) = \theta(t) + \eta e \quad (2.7)$$

where:

- $w_i(t) \in \mathbb{R}$ denotes the weight associated with feature x_i at iteration t ,
- $w_i(t+1)$ denotes the updated weight after the learning step,

- $x_i \in \mathbb{R}$ represents the i -th input feature,
- $\theta(t) \in \mathbb{R}$ denotes the bias term at iteration t ,
- $\theta(t+1)$ denotes the updated bias term,
- $\eta > 0$ is the learning rate controlling the magnitude of parameter updates,
- e is the prediction error defined in Equation 2.5,
- $t \in \mathbb{N}$ denotes the iteration index.

The learning process involves iteratively selecting training samples, computing predictions, evaluating errors, and updating parameters. The algorithm converges when the data are linearly separable (Novikoff, 1962; Rosenblatt, 1958). For non-linearly separable data, more advanced models are required.

2.1.5 Example Calculation with Perceptron Learning

To illustrate the learning process, consider a simple dataset with three input features: age, body mass index (BMI), and cholesterol level. The output variable represents the presence of cardiovascular disease (CVD) as a binary label.

Table 2.1: Sample data for perceptron training

Age	BMI	Cholesterol Level	CVD Risk (Desired Output)
40	23	160	1
55	27	200	1
30	22	170	0

Consider the third training sample, where

$$\mathbf{x} = [30, 22, 170]^T$$

and

$$y_{\text{true}} = 0.$$

Let the initial weights be

$$w_{\text{age}} = 0.2, \quad w_{\text{BMI}} = 0.4, \quad w_{\text{chol}} = 0.1$$

and the bias term be

$$\theta = 0.3,$$

with learning rate

$$\eta = 0.1.$$

The weighted input is computed as:

$$s = (0.2 \times 30) + (0.4 \times 22) + (0.1 \times 170) + 0.3$$

$$s = 6 + 8.8 + 17 + 0.3 = 32.1$$

where:

- $s \in \mathbb{R}$ represents the weighted input to the activation function,
- 30, 22, and 170 correspond to the values of age, BMI, and cholesterol level, respectively.

Applying the activation function defined in Equation 2.2:

$$\hat{y} = 1$$

The predicted output indicates the presence of CVD. However, the true target value is

$$y_{\text{true}} = 0,$$

which means that the sample is incorrectly classified. The prediction error is therefore

$$e = y_{\text{true}} - \hat{y} = 0 - 1 = -1.$$

Using Equations 2.6 and 2.7, the updated parameters become:

$$w_{\text{age}} = 0.2 + (0.1)(-1)(30) = -2.8$$

$$w_{\text{BMI}} = 0.4 + (0.1)(-1)(22) = -1.8$$

$$w_{\text{chol}} = 0.1 + (0.1)(-1)(170) = -16.9$$

$$\theta = 0.3 + (0.1)(-1) = 0.2$$

After the parameter update, the model adjusts its weights and bias in a direction that reduces the likelihood of producing the same classification error for similar observations. The updated weights decrease the influence of the input variables that contributed to the incorrect prediction, while the bias adjustment shifts the overall position of the decision boundary. The perceptron performs classification using a linear decision function of the form

$$w_1x_1 + w_2x_2 + \dots + w_nx_n + \theta = 0,$$

which defines a linear boundary that separates the feature space into two output classes. In this example, age, BMI, and cholesterol level are combined linearly through their

corresponding weights to determine whether a patient is classified as having CVD or not. Repeated updates across multiple training samples gradually modify the position of this boundary and improve classification performance. For linearly separable datasets, the perceptron learning algorithm can converge to a stable solution that correctly separates the classes (Rosenblatt, 1958). However, real-world cardiovascular datasets are often more complex and may contain nonlinear relationships between clinical variables. For example, the combined influence of age, blood pressure, cholesterol level, and lifestyle-related factors may vary across patient groups and cannot always be represented effectively using a single linear boundary. This limitation reduces the ability of the perceptron to model complex clinical patterns accurately. As a result, more advanced neural network models, such as Feedforward Neural Networks (FFNNs), are required. By introducing hidden layers and nonlinear activation functions, FFNNs can learn more flexible decision boundaries and capture more complex relationships within clinical data.

2.1.6 Perceptron Algorithm Training Process

Extending the example above, the perceptron training process involves iterative updates across the full dataset. For each training sample, the algorithm computes a predicted output, evaluates the prediction error, and updates the model parameters using the perceptron learning rule (Haykin, 1999). This process is repeated sequentially across all training samples over multiple epochs until convergence or a predefined stopping criterion is reached. The main components involved in the training process include:

- $w_i \in \mathbb{R}$: weight associated with the i -th input feature,
- $\theta \in \mathbb{R}$: bias term,
- $\eta > 0$: learning rate controlling the update magnitude,
- $x_i \in \mathbb{R}$: input feature value,
- $e \in \mathbb{R}$: prediction error,
- $t \in \mathbb{N}$: iteration or training step index.

The objective of training is to minimise classification errors by progressively refining the decision boundary. In this study, the perceptron is introduced as a foundational learning model that provides the theoretical basis for understanding more advanced neural network architectures. Although the final comparative analysis in this thesis focuses on the Feedforward Neural Network (FFNN) baseline model and the proposed hybrid framework, understanding the perceptron is important because FFNNs are constructed from multiple interconnected perceptron-like units. The perceptron is applied here to clinical variables such as age, BMI, and cholesterol level to demonstrate the fundamental principles of supervised learning, weight updates, and binary classification. While the example demonstrates a small feature set, the same approach extends naturally to higher-dimensional datasets. Input variables are normalised before training to improve numerical stability

and learning efficiency. However, the perceptron remains a linear classifier and performs effectively only when the data are linearly separable. In real-world clinical datasets, relationships between variables are often nonlinear and more complex. This limitation motivates the use of more advanced models, such as multilayer neural networks, which can capture nonlinear patterns and more sophisticated decision boundaries (Goodfellow et al., 2016; Rumelhart et al., 1986). This naturally leads to the use of Feedforward Neural Networks, which are introduced in the following section as the primary modelling approach and baseline comparison model for this study.

2.2 Feedforward Neural Networks for Cardiovascular Risk Prediction

Feedforward Neural Networks (FFNNs) extend the perceptron model introduced in Section 2.1.3. While the perceptron defines a linear mapping between input features and output (see Equation 2.1), FFNNs introduce multiple layers and nonlinear activation functions. This enables the model to capture more complex relationships in the data (Bishop, 2006; Goodfellow et al., 2016). An FFNN consists of an input layer, one or more hidden layers, and an output layer. Each neuron computes a weighted sum of its inputs and applies an activation function, extending the single-layer formulation of the perceptron to a multi-layer structure. Formally, for a neuron j in layer l , the output activation is given by:

$$a_j^{(l)} = \phi \left(\sum_i w_{ji}^{(l)} a_i^{(l-1)} + b_j^{(l)} \right)$$

where:

- $a_j^{(l)} \in \mathbb{R}$ denotes the output activation of neuron j in layer l ,
- $a_i^{(l-1)} \in \mathbb{R}$ denotes the input activation received from neuron i in the previous layer $(l - 1)$,
- $w_{ji}^{(l)} \in \mathbb{R}$ is the weight connecting neuron i in layer $(l - 1)$ to neuron j in layer l ,
- $b_j^{(l)} \in \mathbb{R}$ is the bias term associated with neuron j ,
- $\phi(\cdot)$ denotes the activation function,
- $l \in \mathbb{N}$ represents the layer index.

The quantity

$$z_j^{(l)} = \sum_i w_{ji}^{(l)} a_i^{(l-1)} + b_j^{(l)}$$

represents the weighted input (pre-activation value) for neuron j in layer l . Common activation functions include sigmoid, ReLU, and tanh, which allow the network to model nonlinear patterns commonly observed in clinical data.

In this study, the FFNN is implemented as a supervised learning model for predicting cardiovascular disease (CVD) risk. The network takes patient-level clinical features as input and produces a probability of disease occurrence at the output layer. Unlike the hybrid framework introduced in Chapter 3, this model is trained on the full dataset without subgroup separation. As a result, it represents a global prediction approach and serves as a baseline model for evaluating the advantages of the proposed hybrid framework. The following subsections describe the key components of the FFNN model, including activation functions, parameter learning, and network architecture.

2.2.1 Activation Function and Output Node

Each neuron in the network computes a weighted sum of its inputs followed by an activation function. In this study, the Rectified Linear Unit (ReLU) activation function is used in all hidden layers and is defined as:

$$\text{ReLU}(z) = \max(0, z)$$

where:

- $z \in \mathbb{R}$ denotes the weighted input (pre-activation value),
- $\text{ReLU}(z) \in \mathbb{R}_{\geq 0}$ denotes the output of the activation function.

ReLU introduces nonlinearity while remaining computationally efficient and less prone to vanishing gradient problems compared to traditional activation functions (Goodfellow et al., 2016; Nair & Hinton, 2010). The output layer consists of a single neuron with a sigmoid activation function defined as:

$$\sigma(z) = \frac{1}{1 + \exp(-z)}$$

where:

- $\sigma(z)$ denotes the sigmoid activation output,
- $\sigma(z) \in (0, 1)$ represents the predicted probability produced by the sigmoid function,
- $z \in \mathbb{R}$ is the weighted input to the output neuron.

The sigmoid activation produces a probability estimate $\hat{y} \in (0, 1)$, where $\hat{y}$ represents the estimated probability of cardiovascular disease occurrence. A threshold of 0.5 is applied to obtain the final class label, where values greater than or equal to 0.5 indicate the presence of CVD.

2.2.2 Weights and Biases

Weights and biases are the learnable parameters of the FFNN. Weights determine the contribution of each input feature or hidden-layer activation to the next layer, while biases

allow the activation values to shift and improve the flexibility of the model. In this study, these parameters are learned automatically during training using backpropagation, which extends the basic weight-update principle of the perceptron to multi-layer neural networks (Rumelhart et al., 1986). During training, the network compares the predicted output with the true CVD label using the Binary Cross-Entropy loss function defined in Section 2.2.3. The resulting error is propagated backward through the network, and the weights and biases are updated iteratively using the Adam optimisation algorithm. This process enables the FFNN to reduce prediction error and learn nonlinear relationships between the nine clinical input features and the binary CVD outcome. In this implementation, the learned weights and biases are not interpreted individually. Instead, they are treated as internal model parameters that enable the FFNN to generate predicted CVD probabilities for unseen patients. The predictive performance of the trained network is then evaluated using the metrics described in Section 2.2.5.

2.2.3 Model Architecture and Training Setup

Each patient is represented by a feature vector $\mathbf{x} \in \mathbb{R}^9$, as described in Section 1.3, where the nine dimensions correspond to demographic, clinical, and lifestyle variables associated with cardiovascular risk. Continuous variables are normalised using min–max scaling to ensure numerical stability and improve convergence during training, while categorical variables are encoded numerically for direct processing by the network.

Data Partitioning The dataset is divided into training and testing subsets using an 80:20 split. Stratified sampling is applied to preserve the class distribution of CVD and non-CVD cases in both subsets, ensuring a reliable and unbiased evaluation of model performance. The FFNN adopts a fully connected structure consisting of an input layer, multiple hidden layers, and an output layer. In this study, the network depth is fixed to three hidden layers, while the number of neurons in each layer is varied across configurations to evaluate model capacity. The general architecture is defined as follows:

- Input layer with 9 neurons
- Three hidden layers with varying numbers of neurons (ReLU activation)
- Output layer with 1 neuron (sigmoid activation)

Fixing the number of hidden layers enables a controlled comparison of model capacity by varying only the number of neurons, while maintaining a consistent structure. A moderate depth of three hidden layers is selected to balance expressive power and generalisation. Although deeper networks can model highly complex relationships, they are not always necessary for structured tabular data and may increase the risk of overfitting. The selected architecture therefore provides sufficient flexibility while maintaining stable learning behaviour (Bishop, 2006; Goodfellow et al., 2016). The best-performing configuration is selected based on evaluation results on unseen test data (see Section 2.2.6).

Since the FFNN performs binary classification for CVD prediction, the model outputs a probability value between 0 and 1 using the sigmoid activation function described in Section 2.2.1. To train the network, a loss function is required to measure the difference between the predicted probabilities and the true class labels. In this study, Binary Cross-Entropy (BCE) loss is used because it is specifically designed for binary classification problems and is well suited for probabilistic outputs produced by sigmoid activation functions (Goodfellow et al., 2016). The BCE loss function is defined as:

$$\mathcal{L} = -\frac{1}{N} \sum_{i=1}^N [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$$

where:

- $\mathcal{L} \in \mathbb{R}_{\geq 0}$ denotes the Binary Cross-Entropy loss,
- $N \in \mathbb{N}$ is the number of training samples,
- $y_i \in \{0, 1\}$ represents the true class label for sample i ,
- $\hat{y}_i \in (0, 1)$ denotes the predicted probability produced by the network for sample i ,
- $\log(\cdot)$ denotes the natural logarithm function.

The BCE loss function penalises incorrect predictions more strongly when the predicted probability differs substantially from the true class label. Minimising this loss enables the network to improve prediction accuracy during training.

Parameter optimisation is performed using the Adam algorithm with a learning rate of 0.001 (Kingma & Ba, 2015), enabling efficient and stable convergence during training. The training process is configured as follows:

- Learning rate: 0.001
- Batch size: 64
- Maximum number of epochs: 200
- Early stopping patience: 10 epochs based on validation loss

A validation split of 20% of the training data is used to monitor performance during training. Early stopping is applied to prevent overfitting by terminating training when validation loss no longer improves (Prechelt, 1998). The best model weights are restored automatically to ensure optimal generalisation. This configuration provides a stable and efficient training process, forming a strong baseline for comparison with the hybrid modelling approach introduced in the next chapter.

2.2.4 Model Selection and Architecture Visualisation

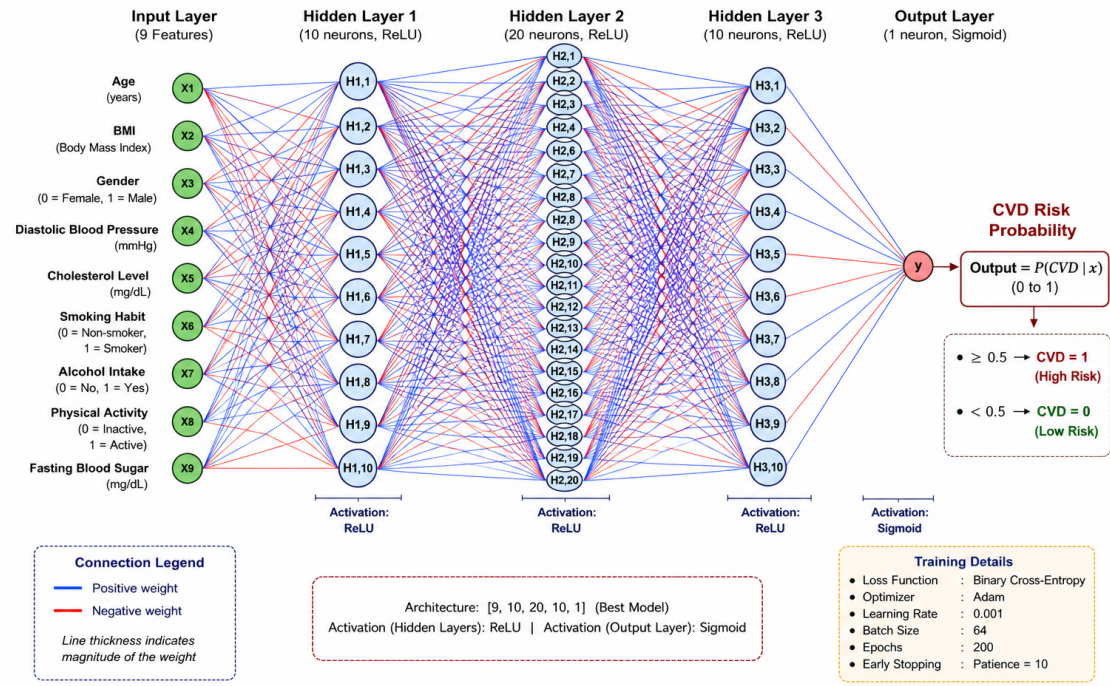


Figure: Feedforward Neural Network Architecture for Cardiovascular Risk Prediction.
The model takes 9 clinical and lifestyle features as input and outputs the probability of cardiovascular disease (CVD).

Figure 2.3: Best-performing FFNN architecture selected based on experimental evaluation. The network consists of nine input features, three hidden layers with ReLU activation, and a single output neuron with sigmoid activation for CVD probability estimation.

To determine the most suitable network structure, a set of candidate architectures is evaluated by varying the number of neurons in each hidden layer while keeping the network depth fixed (see Section 2.2). In total, ten architectures are considered, providing a structured and computationally efficient exploration of model capacity. The architectures follow a systematic pattern in which the number of neurons increases progressively across configurations, allowing analysis of how model complexity influences performance. In addition, symmetric layer structures (e.g., 10–20–10 and 20–50–20) are used to maintain balanced information flow across layers, which is a common design principle in fully connected neural networks (Goodfellow et al., 2016). Each configuration is trained under identical conditions and evaluated on the unseen test dataset (see Section 2.2.3). The architecture achieving the highest classification accuracy is selected as the final model. The selected architecture is visualised in Figure 2.3.

2.2.5 Evaluation Metrics

Model performance is evaluated primarily using classification accuracy. However, in clinical prediction tasks, relying on a single evaluation metric may provide an incomplete assessment of model behaviour and practical usefulness. Therefore, additional evaluation

measures are considered to provide a more comprehensive analysis of predictive performance (Sokolova & Lapalme, 2009). These metrics are derived from the confusion matrix shown in Figure 2.5, which summarises the relationship between predicted class labels and the true clinical outcomes. The confusion matrix provides the numbers of correctly and incorrectly classified observations and forms the basis for evaluating different aspects of classification performance.

- **Accuracy:** overall proportion of correctly classified observations, defined as

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

- **Precision:** reliability of positive predictions, defined as

$$\text{Precision} = \frac{TP}{TP + FP}$$

- **Recall:** ability to correctly identify CVD cases, defined as

$$\text{Recall} = \frac{TP}{TP + FN}$$

- **F1-score:** harmonic mean of precision and recall, defined as

$$\text{F1-score} = \frac{2 \cdot (\text{Precision} \cdot \text{Recall})}{\text{Precision} + \text{Recall}}$$

- **AUC–ROC:** area under the Receiver Operating Characteristic (ROC) curve, measuring the ability of the model to distinguish between classes across different decision thresholds.

where:

- $TP \in \mathbb{N}$ denotes the number of true positive predictions,
- $TN \in \mathbb{N}$ denotes the number of true negative predictions,
- $FP \in \mathbb{N}$ denotes the number of false positive predictions,
- $FN \in \mathbb{N}$ denotes the number of false negative predictions.

Together, these metrics provide a balanced evaluation of predictive performance and clinical relevance, particularly in datasets where class distributions and decision thresholds may vary.

2.2.6 Architecture Performance Analysis

The performance of all evaluated architectures is summarised in Table 2.2.

Table 2.2: Performance comparison of FFNN architectures (best model highlighted)

Architecture	Accuracy	Precision	Recall	F1-score	AUC
[10, 20, 10, 1]	0.7445	0.7380	0.7562	0.7470	0.7462
[15, 30, 15, 1]	0.7243	0.7112	0.7530	0.7315	0.7468
[20, 50, 20, 1]	0.7234	0.7046	0.7670	0.7345	0.7474
[25, 60, 25, 1]	0.7309	0.7239	0.7445	0.7340	0.7489
[30, 70, 30, 1]	0.7296	0.7225	0.7433	0.7327	0.7483
[35, 80, 35, 1]	0.7297	0.7214	0.7463	0.7336	0.7487
[40, 90, 40, 1]	0.7301	0.7277	0.7332	0.7304	0.7484
[45, 100, 45, 1]	0.7285	0.7200	0.7455	0.7326	0.7485
[50, 110, 50, 1]	0.7308	0.7279	0.7349	0.7314	0.7488
[60, 120, 60, 1]	0.7315	0.7285	0.7359	0.7322	0.7486

The best-performing model is highlighted in bold. The results indicate that the architecture [10, 20, 10, 1] achieves the highest classification accuracy (74.45%) and is therefore selected as the final model. A clear pattern can be observed across configurations. After the initial model, accuracy decreases slightly and then stabilises within a narrow range of approximately 72%–73%. This suggests that increasing the number of neurons does not lead to consistent improvements in predictive performance. In contrast, AUC values remain relatively stable across all architectures, indicating that the overall discriminative ability of the models is largely unaffected by increased complexity. However, this stability does not translate into improved classification outcomes. From a modelling perspective, this reflects diminishing returns as network size increases. Larger architectures introduce more parameters but do not provide meaningful gains in predictive accuracy. These findings suggest that a relatively simple network structure is sufficient to capture the main patterns in the dataset. Increasing complexity beyond this point may introduce unnecessary model capacity without improving generalisation. This behaviour is commonly observed in structured tabular data, where overly complex neural networks do not always outperform simpler models. Overall, the selected architecture provides a strong and efficient baseline for cardiovascular risk prediction. It balances predictive performance and model simplicity, which is important for reliable evaluation and comparison with the hybrid framework presented in the following chapter.

Figure 2.4 illustrates the variation in classification accuracy across different FFNN architectures.

2.2.7 Accuracy Trend Across Architectures

The results show that the highest accuracy is achieved by the first architecture, followed by a noticeable decline. After this initial drop, accuracy values stabilise, forming a plateau across subsequent configurations. This pattern indicates diminishing returns as model complexity increases. Increasing the number of neurons does not lead to improved performance and, in some cases, results in a slight reduction in accuracy. This behaviour suggests that the dataset does not benefit from highly complex neural network structures.

Instead, a relatively simple architecture is sufficient to capture the main predictive patterns. This observation is consistent with findings in structured tabular learning, where simpler models often perform competitively compared to more complex deep learning architectures (Shwartz-Ziv & Armon, 2022). Overall, the results highlight the importance of selecting an appropriate level of model complexity rather than assuming that larger networks will necessarily yield better performance.

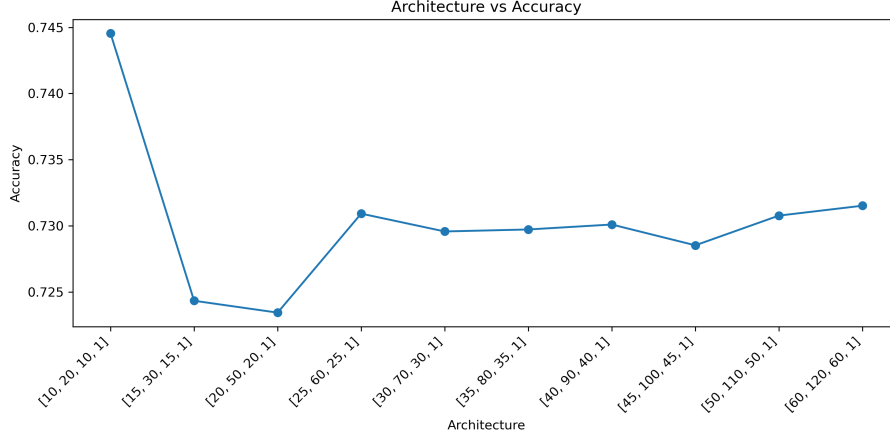


Figure 2.4: Accuracy comparison across FFNN architectures

2.2.8 Best Model Performance

The selected architecture ([10, 20, 10, 1]) achieved an accuracy of **74.45%**, with a precision of **73.80%**, recall of **75.62%**, F1-score of **74.70%**, and an AUC value of **0.7462** on the test dataset. These results indicate balanced classification performance, with no strong bias toward either class. The close alignment between precision and recall suggests stable and consistent prediction behaviour across both positive and negative classes. From a clinical perspective, the slightly higher recall is particularly important, as it improves the identification of patients with cardiovascular disease and reduces the likelihood of missing high-risk cases (Sokolova & Lapalme, 2009). This balance between sensitivity and precision supports the practical applicability of the model in cardiovascular risk prediction settings, where both accurate disease detection and controlled false positive rates are clinically relevant.

2.2.9 Prediction Process and Model Output

After training, the FFNN is used to generate predictions for unseen patient data. Each patient is represented by an input vector

$$\mathbf{x} \in \mathbb{R}^9$$

defined by the clinical variables described in Section 1.3. These inputs are processed through the trained network using a forward pass, following the architecture and learning procedure

outlined in Sections 2.2.1, 2.2.2, and 2.2.3. During this forward pass, the input vector is propagated through each hidden layer, where weighted sums and nonlinear activation functions are applied. At the output layer, a sigmoid activation function produces a probability estimate

$$\hat{y} \in [0, 1]$$

representing the predicted likelihood of cardiovascular disease (CVD), where values closer to 1 indicate higher predicted risk and values closer to 0 indicate lower predicted risk. This probabilistic output provides a continuous measure of risk rather than a fixed categorical decision. To obtain a final classification, a decision threshold

$$\tau = 0.5$$

is applied:

$$\hat{y}_{\text{class}} = \begin{cases} 1, & \text{if } \hat{y} \geq \tau \\ 0, & \text{if } \hat{y} < \tau \end{cases}$$

where:

- $\hat{y}_{\text{class}} \in \{0, 1\}$ denotes the predicted class label,
- $\hat{y}$ denotes the predicted probability generated by the sigmoid output node,
- $\tau \in [0, 1]$ denotes the classification threshold.

A value of 1 indicates the presence of CVD and a value of 0 indicates its absence. This threshold-based decision converts the predicted probability into a clinically interpretable binary outcome. The prediction process is applied to all samples in the test dataset introduced in Section 2.2.3. The predicted labels are then compared with the true outcomes to evaluate model performance using the metrics described in Section 2.2.5. This formulation ensures that the FFNN produces both probabilistic risk estimates and discrete classifications, supporting quantitative evaluation as well as potential clinical interpretation of patient risk.

2.2.10 Results and Performance Analysis

Confusion Matrix Analysis

Using the evaluation metrics defined in Section 2.2.5, the classification behaviour of the FFNN model is further analysed through the confusion matrix presented in Figure 2.5. The confusion matrix summarises the agreement between predicted class labels and the true clinical outcomes, allowing the identification of correctly and incorrectly classified observations.

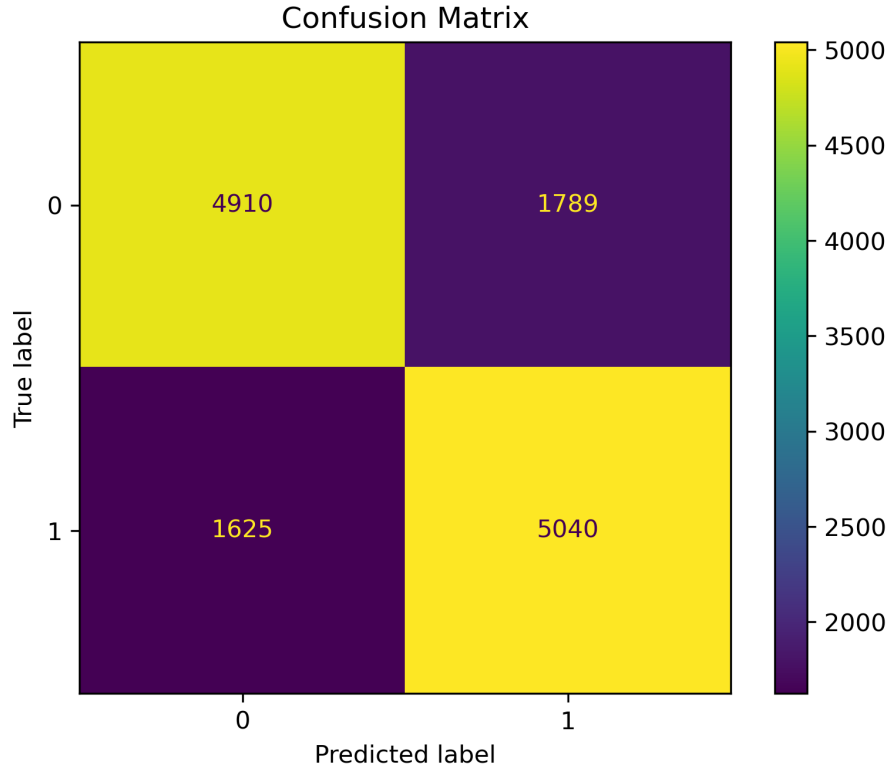


Figure 2.5: Confusion matrix for the best FFNN model

The model correctly classified **5040** CVD cases as positive and **4910** non-CVD cases as negative. In addition, **1789** non-CVD cases were incorrectly classified as CVD, while **1625** CVD cases were incorrectly classified as non-CVD. These results indicate that the FFNN achieves strong overall classification performance while maintaining relatively balanced behaviour across both classes. From a clinical perspective, false negative predictions are particularly important because they correspond to patients with cardiovascular disease who are incorrectly identified as low risk. Such errors may delay clinical intervention or preventive treatment (Harrell, 2015). False positive predictions, although generally less critical clinically, may still increase unnecessary follow-up examinations or additional diagnostic procedures. Overall, the confusion matrix demonstrates that the FFNN is capable of identifying meaningful patterns in the clinical data while maintaining balanced sensitivity and specificity characteristics.

Loss Curve Analysis

The training and validation loss curves shown in Figure 2.6 illustrate the learning behaviour of the model during training. The loss values represent the Binary Cross-Entropy (BCE) loss defined in Section 2.2.3, where lower values indicate better agreement between predicted probabilities and true class labels. Both curves demonstrate smooth convergence, with a relatively small gap between training and validation loss. This indicates stable optimisation and good generalisation performance. The model converges within a limited number of

epochs, suggesting efficient parameter learning. The similar trend between training and validation loss suggests that overfitting is minimal. This behaviour is consistent with well-regularised neural networks trained using early stopping (Prechelt, 1998).

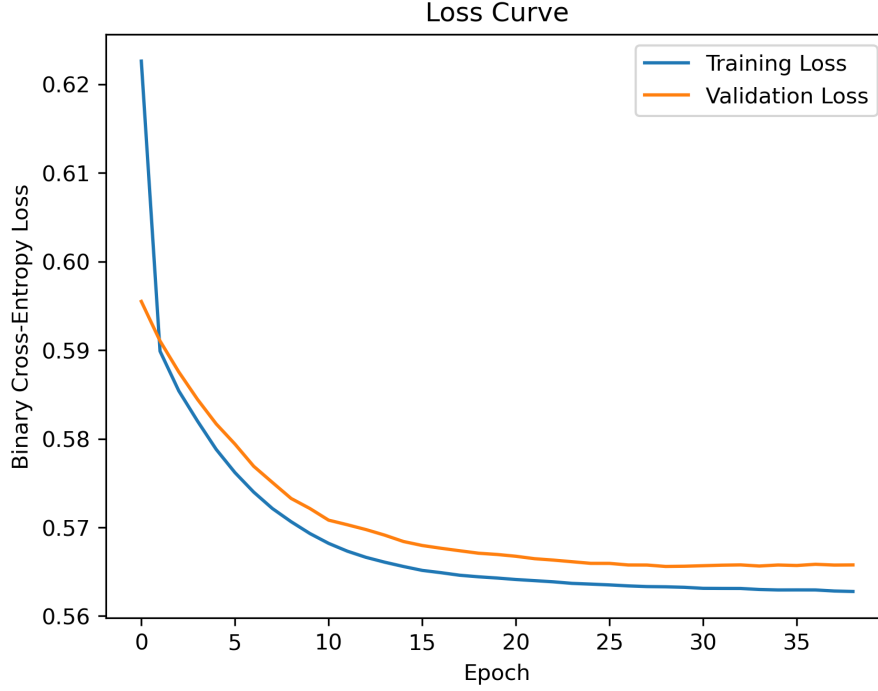


Figure 2.6: Training and validation loss curves

ROC Curve Analysis

The ROC curve in Figure 2.7 evaluates the discriminative performance of the model based on the predicted probabilities generated in the prediction step (Section 2.2.9). The Receiver Operating Characteristic (ROC) curve represents the relationship between the true positive rate and the false positive rate across different classification thresholds.

The true positive rate is defined as:

$$TPR = \frac{TP}{TP + FN}$$

and the false positive rate is defined as:

$$FPR = \frac{FP}{FP + TN}$$

where TP , TN , FP , and FN have the meanings defined in Section 2.2.10.

The model achieves an Area Under the Curve (AUC) value of approximately 0.7462, indicating moderate ability to distinguish between CVD and non-CVD cases. The gradual shape of the curve suggests that the separation between classes is not strong. This indicates that the model has limited ability to clearly differentiate between high-risk and low-risk patients using the current feature set and architecture (Fawcett, 2006). As a result,

prediction uncertainty remains relatively high across certain threshold ranges.

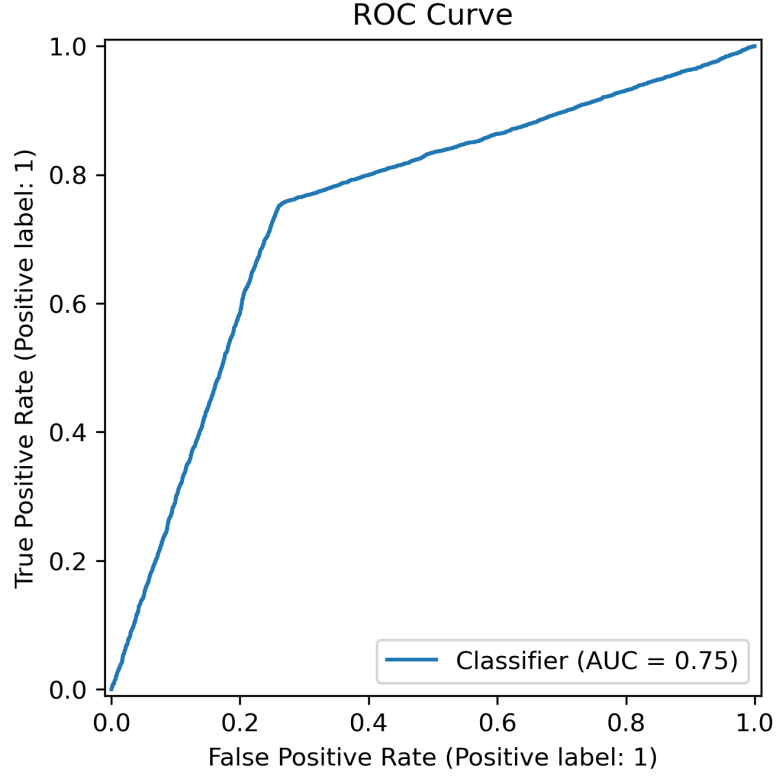


Figure 2.7: ROC curve for the best FFNN model

Summary of Results

The FFNN model provides a baseline for cardiovascular disease risk prediction. The reported results are based on predictions generated on the unseen test dataset using the trained model (see Section 2.2.9). The selected architecture achieves consistent performance across evaluation metrics (see Section 2.2.8). The analysis shows that increasing model complexity does not lead to improved accuracy (Section 2.2.6), with performance remaining relatively stable across different configurations. The model demonstrates moderate discriminative ability (Section 2.2.10) and produces a non-negligible number of misclassified cases, particularly false negatives (Section 2.2.10). Overall, these findings highlight both the strengths and limitations of the FFNN approach. While the model is capable of learning nonlinear relationships and achieving reasonable predictive performance, its limitations in discrimination and error patterns motivate the development of a more structured and interpretable framework. These limitations form the basis for the hybrid modelling approach introduced in Chapter 3.

2.3 Discussion

The results obtained using the CVD SL dataset (see Section 1.3) indicate that the FFNN model is able to learn relevant patterns from the clinical data and produce stable predictions.

However, the observed performance trends suggest that increasing model capacity alone is not sufficient to improve predictive outcomes in this context. The limited impact of increasing architectural complexity (Section 2.2.6) reflects the structured nature of the dataset, where additional parameters do not translate into improved generalisation. This behaviour is consistent with findings in prior studies on neural networks applied to tabular clinical data (Bishop, 2006; Detrano et al., 1989; Goodfellow et al., 2016; Krittanawong et al., 2020). The training process demonstrates stable convergence and effective optimisation (Section 2.2.10), indicating that the model is appropriately fitted under the selected configuration. However, the primary limitations arise from the underlying model structure. As a global model, the FFNN represents the entire population using a single decision function, which restricts its ability to capture variation in risk patterns across different patient profiles. This limitation is particularly important in clinical datasets, where patient heterogeneity plays a central role in disease progression and risk assessment. In addition, the lack of interpretability limits the practical applicability of the model in healthcare settings, where understanding the reasoning behind predictions is essential for clinical trust and decision-making (Rudin, 2019). The absence of temporal modelling further reduces its ability to represent how cardiovascular risk evolves over time, which is a key aspect of chronic disease progression. Taken together, these findings suggest that improvements in cardiovascular risk prediction require changes in modelling structure rather than further increases in model complexity. Specifically, there is a need for approaches that can account for patient heterogeneity, provide interpretable outputs, and represent risk as a time-dependent process. These requirements directly motivate the development of the structured hybrid framework introduced in Chapter 3.

2.4 Conclusion

This chapter presented an evaluation of Feedforward Neural Networks (FFNNs) for cardiovascular disease risk prediction using the CVD SL dataset, establishing a reliable baseline model. The results demonstrate that the FFNN is capable of learning nonlinear relationships and producing stable predictions, providing a useful reference for comparison. However, several limitations are identified. The model operates as a global predictor, limiting its ability to capture variation across patient subgroups. Its lack of interpretability restricts clinical usability, and the absence of temporal modelling prevents meaningful representation of risk progression over time. These limitations highlight the challenges of applying standard neural network approaches in clinical risk prediction tasks. The findings of this chapter therefore provide a clear motivation for the development of more structured and interpretable modelling approaches. In particular, they emphasise the need to incorporate subgroup-level analysis and time-dependent risk modelling. The next chapter introduces a hybrid computational framework that integrates clustering, survival analysis, and cumulative risk estimation to address these limitations and provide a more clinically meaningful representation of cardiovascular risk.

CHAPTER 3

A HYBRID COMPUTATIONAL FRAMEWORK FOR CARDIOVASCULAR RISK STRATIFICATION

3.1 Introduction

Accurate and interpretable prediction of cardiovascular disease (CVD) risk remains a major challenge in modern healthcare analytics. As discussed in Chapter 2, conventional global prediction models can learn important patterns from clinical data, but they often have limited ability to represent heterogeneity across patient populations. Traditional statistical approaches, particularly regression-based models, are widely used in clinical practice because of their simplicity and interpretability. However, these methods rely on strong assumptions, including linear relationships between variables and population homogeneity. In real-world clinical settings, these assumptions are frequently violated, since patient populations are heterogeneous and cardiovascular risk factors may interact in complex and potentially nonlinear ways (Harrell, 2015; Hastie et al., 2009). Consequently, global models may fail to capture subgroup-specific risk patterns and may not provide sufficiently personalised risk estimates. The FFNN baseline model presented in Chapter 2 addressed some limitations of traditional statistical approaches by learning nonlinear relationships directly from clinical variables. However, the results also demonstrated several important limitations. In particular, the FFNN operated as a single global model and therefore could not explicitly represent variation in cardiovascular risk across different patient subgroups. In addition, the lack of interpretability and the absence of temporal risk modelling reduced its clinical applicability. These observations motivate the need for a more structured modelling framework that combines predictive flexibility with clinically meaningful interpretation.

With the increasing availability of large-scale clinical datasets, there is a growing need for computational approaches that can model complex data structures while remaining interpretable and clinically useful. In particular, effective cardiovascular risk modelling should account for heterogeneity within patient populations, where different groups of patients may exhibit distinct disease progression patterns and risk profiles (Molnar, 2022; Rudin, 2019). This motivates the development of hybrid methodologies that combine data-driven learning with survival-based risk modelling strategies. In this context, this chapter proposes a hybrid computational framework for cardiovascular risk stratification that integrates subgroup discovery, survival analysis, and cumulative risk estimation (see Section 3.4). The central idea is to identify clinically meaningful patient subgroups directly from data and then model cardiovascular risk separately within each subgroup. Unlike the FFNN baseline introduced in Chapter 2, the proposed framework does not treat the entire population as a single homogeneous group. Instead, it combines clustering and survival modelling to support subgroup-specific and time-dependent risk estimation.

Key Contributions of This Chapter This chapter makes three primary contributions:

- A subgroup-based survival modelling framework that integrates Winner-Takes-All (WTA) competitive learning with cluster-specific Cox proportional hazards models, enabling modelling of heterogeneous patient populations.

- The introduction of the Cumulative Prevalence Ratio (CPR) as an interpretable measure of long-term cardiovascular risk, summarising time-dependent hazard information into a single comparable value across patient subgroups.
- A unified computational pipeline that combines clustering, survival analysis, and risk stratification to support both predictive performance and clinical interpretability.

A key component of the proposed framework is the Winner-Takes-All (WTA) competitive learning mechanism (Haykin, 1999; Kröse & van der Smagt, 1996). WTA is a form of competitive learning in which multiple computational units compete to represent an input pattern, and only the unit with the highest activation, referred to as the “winner”, is selected and updated (Kroese & van der Smagt, 1996). Through this process, the model learns prototype vectors that capture dominant structures and similarities within the dataset. In the proposed framework, WTA is used to partition the dataset into patient subgroups with similar clinical characteristics by assigning each patient to the most representative prototype unit (see Section 3.4.1). Following subgroup identification, cardiovascular risk is modelled separately within each subgroup using cluster-specific Cox proportional hazards models (Cox, 1972). This allows the relationship between predictors and cardiovascular outcomes to vary across patient groups, improving modelling flexibility and supporting more personalised risk estimation (see Section 3.4.2).

To support long-term cardiovascular risk assessment, the framework further introduces the Cumulative Prevalence Ratio (CPR), which provides a summary measure of cumulative risk over a specified time interval (see Section 3.4.3). Unlike single-point probability estimates, CPR captures the progression of risk over time and enables direct comparison of cumulative cardiovascular burden across patient subgroups. This supports more clinically meaningful interpretation and risk stratification. The overall workflow of the proposed framework consists of four main stages:

1. Subgroup identification using WTA competitive learning,
2. Estimation of cluster-specific Cox survival models,
3. Computation of cumulative cardiovascular risk using CPR,
4. Assignment of patients to clinically meaningful risk categories.

These stages are implemented within a unified modelling pipeline (see Algorithm 1), providing a structured, interpretable, and time-dependent approach to cardiovascular risk stratification. The framework is particularly suitable for large and heterogeneous clinical datasets and supports personalised risk assessment by combining subgroup discovery with survival-based modelling. The following sections describe each framework component in detail.

3.2 Related Work

Cardiovascular disease (CVD) remains one of the leading causes of morbidity and mortality worldwide, motivating extensive research into reliable and interpretable risk prediction models (Benjamin et al., 2017; World Health Organization, 2021). Traditional statistical approaches, particularly logistic regression and the Cox proportional hazards model, have formed the foundation of clinical risk assessment for several decades (Cox, 1972; Therneau & Grambsch, 2000). These models provide interpretable relationships between risk factors and clinical outcomes and have been widely applied in large cohort studies. Established tools such as the Framingham Risk Score and QRISK3 are based on these principles and remain widely used in clinical practice because of their transparency and ease of interpretation (Hippisley-Cox & Coupland, 2017; Wilson et al., 1998). Despite their advantages, classical statistical models rely on simplifying assumptions, including linear relationships between predictors and outcomes, proportional effects over time, and population homogeneity. In real-world clinical datasets, these assumptions are often violated, since patient populations are heterogeneous and cardiovascular risk factors may interact in complex and potentially nonlinear ways. Consequently, global models may not provide sufficiently accurate personalised risk estimates, particularly for patients with distinct clinical characteristics or atypical risk profiles. To address these limitations, machine learning approaches have increasingly been applied to cardiovascular risk prediction (Krittanawong et al., 2020; Topol, 2019). Methods such as random forests, gradient boosting, support vector machines, and neural networks are capable of modelling nonlinear relationships and complex feature interactions without requiring strict statistical assumptions (Chen & Guestrin, 2016). In particular, neural network models such as Feedforward Neural Networks (FFNNs) can learn high-dimensional representations directly from data and may improve predictive performance in complex clinical settings. However, although these approaches often achieve strong classification performance, they frequently lack interpretability and transparency. This limits their direct applicability in healthcare environments, where clinicians require understandable reasoning behind model predictions to support trust and decision-making (Rudin, 2019). As a result, a fundamental trade-off remains between predictive performance and interpretability.

Another important research direction focuses on clustering and subgroup discovery. Unsupervised learning methods enable the identification of latent patient subgroups that share similar clinical characteristics or disease patterns. Such approaches are particularly relevant in cardiovascular medicine, where heterogeneous populations may exhibit different progression pathways and risk behaviours. Competitive learning methods and self-organising neural networks have been widely used for clustering high-dimensional data (Haykin, 1999; Kröse & van der Smagt, 1996). In competitive learning, multiple computational units compete to represent input patterns, and only the unit with the highest similarity to the input is selected as the winner and updated (Krosé & van der Smagt, 1996). This competitive process enables the formation of prototype-based representations of the

dataset and supports subgroup identification without requiring predefined class labels. However, although these methods are effective for clustering, they often do not incorporate time-to-event modelling, which is essential for cardiovascular risk analysis and disease progression modelling.

A closely related framework was introduced by Liyanage and Lipnicka (Liyanage & Lipnicka, 2026), developed as part of the present doctoral research, and it forms the methodological foundation for this thesis. In that study, patient data were organised into clusters using a centroid-based representation, where each cluster was summarised by the mean feature vector of its members. Following clustering, cluster-specific Cox proportional hazards models were estimated to analyse subgroup-level cardiovascular risk, and survival functions were derived for each subgroup. In addition, the study introduced the Cumulative Prevalence Ratio (CPR) as a measure of long-term cardiovascular risk (see Section 3.4.3). CPR summarises cumulative risk over time and enables comparison between patient subgroups using a single interpretable measure. Although the centroid-based clustering approach is computationally simple and clinically interpretable, it assumes equal contribution of all features during subgroup formation. In practical clinical datasets, however, variables differ substantially in their predictive importance and statistical relevance. Consequently, clustering performance may be influenced by less informative variables, potentially reducing the quality and stability of subgroup identification. To address this limitation, the present thesis introduces a modified framework based on the Winner-Takes-All (WTA) competitive learning mechanism (see Section 3.4.1). In the proposed approach, clustering is performed through a competitive process in which only the most representative computational unit is updated for each input sample. This enables prototype vectors to adapt dynamically to dominant structures within the feature space, allowing feature influence to emerge naturally from the data rather than being treated equally by design.

Based on the above discussion, a clear research gap can be identified. Existing frameworks that combine clustering with survival modelling often rely on approaches that assume equal feature contribution during subgroup formation. Such methods may not adequately represent the relative importance of variables in high-dimensional and heterogeneous clinical datasets. There is therefore a need for computational frameworks that integrate subgroup discovery, feature representation, and time-dependent risk modelling while maintaining interpretability and clinical relevance. The proposed WTA-based framework addresses this gap by embedding competitive learning directly into the clustering process. This supports more robust subgroup identification and improves the reliability of subsequent survival modelling (see Section 3.4.2) and cumulative risk estimation using CPR (see Section 3.4.3). In addition, this chapter includes a comparative analysis between the centroid-based clustering approach and the proposed WTA-based framework, allowing the impact of competitive learning on subgroup formation and cardiovascular risk prediction to be evaluated systematically. Therefore, this thesis contributes to the existing literature by extending cluster-based survival modelling frameworks through the integration of competitive learning, subgroup-specific survival analysis, and cumulative risk estimation

within a unified and interpretable computational pipeline.

3.3 Problem Formulation and Research Objectives

The objective of this research is to develop a computational framework for accurate, interpretable, and clinically meaningful prediction of cardiovascular disease (CVD) risk using patient-level clinical data. The problem is formulated within the framework of survival analysis, which models time-to-event outcomes. In this setting, the aim is not only to determine whether a cardiovascular event occurs, but also to model the timing of its occurrence (Cox, 1972; Hosmer et al., 2008). Let $n \in \mathbb{N}$ denote the total number of patients and $d \in \mathbb{N}$ the number of clinical features, where $d = 9$ in this study. The dataset is defined as

$$\mathcal{D} = \{(x_i, T_i, \delta_i)\}_{i=1}^n, \quad (3.1)$$

where, for each patient $i = 1, \dots, n$:

- $x_i \in \mathbb{R}^d$ denotes the patient feature vector,
- $T_i \in \mathbb{R}^+$ denotes the observed time-to-event or censoring time,
- $\delta_i \in \{0, 1\}$ is the event indicator, where $\delta_i = 1$ indicates event occurrence and $\delta_i = 0$ denotes right censoring.

The feature matrix is denoted by $X \in \mathbb{R}^{n \times d}$, where the i -th row corresponds to the transpose of the patient feature vector, denoted by x_i^T . Here, the superscript T in x_i^T denotes the transpose operator and should not be confused with the survival time variable T_i . The pair (T_i, δ_i) jointly encodes survival time and censoring information. The objective is to model the conditional distribution of the event time given patient features. In general, the proposed framework can be applied to any clinical dataset that contains patient features, an event indicator, and a valid time variable, such as follow-up duration, time from baseline to diagnosis, time to hospitalisation, or time to cardiovascular event. Therefore, the methodology is not limited to age-based modelling and can be adapted to longitudinal datasets where direct follow-up time is available. In the present study, explicit longitudinal follow-up duration was not available in the dataset. Therefore, patient age was used as the observed survival time variable. Accordingly, T_i represents the age of patient i at observation, while δ_i corresponds to the binary cardiovascular disease outcome, where $\delta_i = 1$ indicates the presence of CVD and $\delta_i = 0$ indicates no observed CVD at that age. Under this formulation, the model estimates cardiovascular risk as a function of age and patient-level clinical characteristics. Age is therefore treated as the time scale in the survival model rather than as an ordinary covariate. This formulation allows the model to approximate time-to-event behaviour using cross-sectional data, although it does not represent true longitudinal follow-up. The survival function is defined as

$$S(t | x) = \mathbb{P}(T \geq t | x), \quad (3.2)$$

where $\mathbb{P}(\cdot)$ denotes probability, T is the random variable representing time-to-event, $t \geq 0$ is a specific time point, and $x \in \mathbb{R}^d$ denotes the vector of covariates. The quantity $S(t | x)$ represents the probability of surviving beyond time t given the covariate vector x .

The corresponding hazard function is defined as

$$h(t | x) = \lim_{\Delta t \rightarrow 0} \frac{\mathbb{P}(t \leq T < t + \Delta t | T \geq t, x)}{\Delta t}, \quad (3.3)$$

where $\Delta t > 0$ is a small time interval approaching zero. The quantity $h(t | x)$ describes the instantaneous risk of event occurrence at time t , conditional on survival up to time t .

These quantities are related through

$$S(t | x) = \exp \left(- \int_0^t h(u | x) du \right), \quad (3.4)$$

where $u \in \mathbb{R}^+$ is the integration variable representing time.

To capture heterogeneity in the patient population, the dataset is partitioned into $K \in \mathbb{N}$ subgroups using a Winner-Takes-All (WTA) competitive learning mechanism. WTA is a form of competitive learning in which multiple computational units compete to represent an input pattern, and only the unit with the highest activation is selected as the winner and updated (Kröse & van der Smagt, 1996). For each input vector $x_i \in \mathbb{R}^d$, the activation of cluster $k \in \{1, \dots, K\}$ is defined as

$$s_k = w_k^T x_i, \quad (3.5)$$

where $w_k \in \mathbb{R}^d$ denotes the prototype vector associated with cluster k , $x_i \in \mathbb{R}^d$ is the feature vector of patient i , and $s_k \in \mathbb{R}$ is the corresponding activation score. The cluster assignment is obtained using the WTA rule:

$$c_i = \arg \max_{k \in \{1, \dots, K\}} s_k, \quad (3.6)$$

where $c_i \in \{1, \dots, K\}$ denotes the cluster label assigned to patient i . The quantity s_k measures the similarity between the input vector and the cluster prototype. When both x_i and w_k are normalised, the dot product $w_k^T x_i$ becomes proportional to cosine similarity, and the assignment corresponds to selecting the most similar prototype in the feature space. Following cluster assignment, survival modelling is performed separately within each subgroup. For cluster k , the Cox proportional hazards model is defined as

$$h_k(t | x) = h_{0,k}(t) \exp(\beta_k^T x), \quad (3.7)$$

where $h_k(t | x)$ denotes the hazard function for cluster k , $h_{0,k}(t)$ is the baseline hazard function for cluster k , $\beta_k \in \mathbb{R}^d$ denotes the vector of regression coefficients for cluster k , and $\beta_k^T x \in \mathbb{R}$ represents the linear predictor, also referred to as the log-risk score. This formulation decomposes the global prediction problem into cluster-specific components, allowing the model to capture heterogeneous relationships between covariates and survival

outcomes across different patient subgroups. This approach enables more flexible and personalised risk modelling compared to a single global model. In addition to instantaneous risk estimation, long-term cardiovascular risk is quantified using the Cumulative Prevalence Ratio (CPR), which summarises time-dependent risk over a predefined interval. This provides a complementary measure for comparing cumulative risk across clusters.

The proposed framework is designed to:

- model time-to-event outcomes using survival analysis,
- capture population heterogeneity through clustering,
- estimate subgroup-specific risk using Cox proportional hazards models,
- provide both instantaneous and cumulative measures of cardiovascular risk.

Overall, this formulation establishes a mathematically consistent and conceptually coherent foundation for the proposed hybrid modelling approach.

3.4 Proposed Computational Framework

This section presents the proposed computational framework for cardiovascular disease (CVD) risk stratification. The framework follows a cluster-based survival modelling paradigm, in which patient subgroups are first identified and survival analysis is subsequently performed within each subgroup. A related centroid-based clustering framework was introduced in previous work (Liyanage & Lipnicka, 2026). The present study extends this approach by incorporating a Winner-Takes-All (WTA) competitive learning mechanism for subgroup discovery.

The framework consists of three main stages:

1. Subgroup discovery using WTA-based clustering (see Section 3.4.1),
2. Cluster-specific survival modelling using Cox proportional hazards models (see Section 3.4.2),
3. Cumulative risk estimation using the Cumulative Prevalence Ratio (CPR) (see Section 3.4.3).

This structure enables the modelling of patient heterogeneity while preserving interpretability and clinical relevance.

3.4.1 Stage 1: WTA-Based Clustering

The clustering procedure operates on the dataset defined in Equation (3.1). Each patient is represented by a feature vector $x_i \in \mathbb{R}^d$. Prior to subgroup identification, continuous variables are normalised using min-max scaling to ensure comparable contribution across features, while categorical variables are encoded numerically. Subgroup discovery is

performed using the WTA competitive learning mechanism introduced in Section 3.3. In this study, $K = 15$ prototype vectors were used to represent candidate patient subgroups:

$$(w_k)_{k=1,\dots,K}, \quad w_k \in \mathbb{R}^d$$

where each prototype vector w_k represents the centre of a candidate cluster in the d -dimensional feature space. Each patient is assigned to the cluster whose prototype produces the strongest activation according to the rule defined in Equation (3.6).

The model is trained iteratively. For each input vector, only the winning prototype is updated:

$$w_{k^*} \leftarrow w_{k^*} + \eta(x_i - w_{k^*}), \quad (3.8)$$

where:

- $\eta > 0$ denotes the learning rate,
- $x_i \in \mathbb{R}^d$ is the feature vector of patient i ,
- $w_{k^*} \in \mathbb{R}^d$ is the prototype vector of the winning cluster,
- $k^* \in \{1, \dots, K\}$ denotes the index of the winning unit satisfying

$$k^* = \arg \max_{k \in \{1, \dots, K\}} s_k.$$

During training, the learning rate η decreases gradually across iterations. This allows larger prototype adjustments during the early learning stages and smaller refinements during later stages, helping the clustering process stabilise and converge more smoothly.

The update rule moves the winning prototype towards the input sample and gradually refines subgroup representation within the feature space. As training progresses, prototype vectors adapt to the distribution of patient features, enabling similar patient profiles to be grouped together. The number of clusters $K \in \mathbb{N}$ is a key modelling parameter that controls the granularity of subgroup discovery. Its selection is based on empirical evaluation, considering cluster size distribution, assignment stability, and the reliability of subsequent survival model estimation (see Section 3.4.2). Small values of K produce overly coarse groupings, while large values may result in clusters with insufficient observations. An initial value of $K = 15$ was selected to balance heterogeneity and statistical reliability. Although clustering quality metrics such as the silhouette score were examined, no clear optimal value was identified. This suggests that the dataset does not exhibit a strongly separated clustering structure. Therefore, the choice of K was guided primarily by modelling considerations. Importantly, the competitive learning process naturally suppresses weak or redundant clusters. Prototype units that rarely win receive minimal updates and gradually become inactive. As a result, the effective number of clusters is determined adaptively, producing a stable and parsimonious clustering structure. The resulting cluster assignments are then used to construct subgroup-specific datasets for survival modelling.

3.4.2 Stage 2: Cluster-Specific Survival Modelling

Following subgroup identification (see Section 3.4.1), the dataset is partitioned into cluster-specific subsets:

$$\mathcal{D}_k = \{(x_i, T_i, \delta_i) \mid c_i = k\}, \quad k = 1, \dots, K, \quad (3.9)$$

where $\mathcal{D}_k$ denotes the subset of observations assigned to cluster k , $x_i \in \mathbb{R}^d$ is the feature vector of patient i , $T_i \in \mathbb{R}^+$ denotes the observed time variable, $\delta_i \in \{0, 1\}$ is the event indicator, and c_i is the cluster assignment defined in Equation (3.6). Within each cluster, survival modelling is performed using the Cox proportional hazards formulation introduced in Equation (3.7). The Cox model is selected because of its ability to handle censored data and provide interpretable hazard ratios, which are important in clinical risk analysis.

The hazard function represents the instantaneous risk of a cardiovascular event at time t , conditional on survival up to that time. This formulation relaxes the assumption of a single global relationship between covariates and outcomes. In heterogeneous clinical datasets, the effects of risk factors may vary across patient subgroups. By estimating separate models for each cluster, both the regression coefficients and baseline hazard functions are allowed to vary across subgroups. For each cluster k , the following quantities are estimated:

- Regression coefficient vector $\beta_k \in \mathbb{R}^d$,
- Hazard ratios $\exp(\beta_{k,j})$, where $\beta_{k,j}$ denotes the coefficient associated with feature j ,
- Baseline hazard function $h_{0,k}(t)$,
- Cumulative baseline hazard function $H_{0,k}(t)$.

This provides a flexible and interpretable representation of subgroup-specific cardiovascular risk. The estimated survival models form the basis for cumulative risk computation in the next stage.

3.4.3 Stage 3: Cumulative Risk Estimation using CPR

To quantify long-term cardiovascular risk, the framework employs the Cumulative Prevalence Ratio (CPR), introduced in (Liyanage & Lipnicka, 2026). The CPR aggregates survival risk over a specified time interval $[t_a, t_b]$, producing a scalar risk estimate for each patient subgroup. For cluster k , the CPR is defined as

$$CPR_k = \exp\left(\beta_k^T \bar{x}_k\right) [H_{0,k}(t_b) - H_{0,k}(t_a)], \quad (3.10)$$

where $\beta_k \in \mathbb{R}^d$ denotes the regression coefficient vector estimated from the cluster-specific Cox model, $H_{0,k}(t)$ denotes the cumulative baseline hazard function for cluster k , and $t_a, t_b \in \mathbb{R}^+$ represent the lower and upper bounds of the selected time interval. The vector

$\bar{x}_k \in \mathbb{R}^d$ denotes the mean feature vector of cluster k and is computed as

$$\bar{x}_k = \frac{1}{n_k} \sum_{i:c_i=k} x_i, \quad (3.11)$$

where $x_i \in \mathbb{R}^d$ is the feature vector of patient i , $c_i \in \{1, \dots, K\}$ denotes the cluster assignment obtained from the WTA rule in Equation (3.6), and

$$n_k = |\mathcal{D}_k|$$

represents the number of patients within cluster k . The quantity $\beta_k^T \bar{x}_k$ represents the subgroup-level linear predictor obtained from the inner product between the regression coefficient vector and the mean feature vector. This formulation approximates cumulative subgroup risk by evaluating the Cox model at the average patient profile of each cluster. Higher CPR values indicate greater cumulative cardiovascular risk over the selected time interval, while lower values correspond to more favourable survival profiles. The CPR therefore enables direct comparison of long-term cardiovascular risk across patient subgroups.

3.4.4 Risk Stratification

The final stage of the framework converts continuous CPR estimates into clinically interpretable risk categories. Let

$$\mathbf{CPR} = \{CPR_1, CPR_2, \dots, CPR_K\}$$

denote the set of cluster-specific CPR values defined in Equation (3.10), where $CPR_k \in \mathbb{R}^+$ represents the cumulative cardiovascular risk associated with cluster k . Risk thresholds are determined using percentile-based partitioning:

$$\tau_{0.2} = 20\text{th percentile of } \mathbf{CPR}, \quad (3.12)$$

$$\tau_{0.5} = 50\text{th percentile of } \mathbf{CPR}, \quad (3.13)$$

$$\tau_{0.8} = 80\text{th percentile of } \mathbf{CPR}, \quad (3.14)$$

where $\tau_{0.2}$, $\tau_{0.5}$, and $\tau_{0.8}$ denote the threshold values computed from the empirical CPR distribution.

These thresholds were selected to provide balanced separation between low-, moderate-, high-, and critical-risk groups while maintaining sufficient subgroup size within each category. Each cluster $k \in \{1, \dots, K\}$ is assigned to a risk category according to its CPR value:

- **Low Risk:** $CPR_k \leq \tau_{0.2}$,
- **Moderate Risk:** $\tau_{0.2} < CPR_k \leq \tau_{0.5}$,
- **High Risk:** $\tau_{0.5} < CPR_k \leq \tau_{0.8}$,

- **Critical Risk:** $CPR_k > \tau_{0.8}$.

Each patient inherits the risk category of the cluster assigned through the WTA rule defined in Equation (3.6). This transforms continuous survival-based risk estimates into clinically interpretable categories that support subgroup comparison and targeted cardiovascular risk assessment.

3.4.5 Experimental Setup

Consistent with the problem formulation presented in Section 3.3, age is used as the survival time variable and CVD status is used as the event indicator in the Cox proportional hazards model described in Section 3.4.2. In this setting, each patient is treated as being observed at a specific age, and the model estimates the risk of developing CVD at or before that age as a function of patient-level clinical features. The experimental setup follows the methodology described in Section 3.4 and applies the proposed framework to patient-level clinical data for cardiovascular disease risk analysis. The evaluation is structured around the three main components of the framework, namely WTA-based clustering (Section 3.4.1), cluster-specific survival modelling (Section 3.4.2), and CPR-based cumulative risk estimation (Section 3.4.3). This organisation reflects the integrated nature of the proposed framework, where clustering, survival analysis, and cumulative risk estimation are combined within a unified modelling pipeline. It also ensures that each component is evaluated both independently and in relation to the overall objective of modelling heterogeneous cardiovascular risk. The overall implementation of the proposed framework follows the sequential procedure outlined in Algorithm 1.

Algorithm 1 Hybrid Cardiovascular Risk Stratification Framework

Require: Dataset $\mathcal{D} = \{(x_i, T_i, \delta_i)\}_{i=1}^n$, number of clusters K , learning rate η , number of epochs E , time interval $[t_a, t_b]$

Ensure: Cluster-level CPR values and patient-level risk categories

```

1: Stage 1: Data Preprocessing
2: Clean dataset and handle missing values
3: Encode categorical variables
4: Normalise feature values
5: Stage 2: WTA-Based Subgroup Identification
6: Initialise  $K$  prototype vectors  $w_1, \dots, w_K \in \mathbb{R}^d$ 
7: for epoch = 1 to  $E$  do
8:   for each patient  $x_i$  do
9:     Compute activations  $a_k = w_k^T x_i$  for all  $k$ 
10:     $c_i = \arg \max_{k \in \{1, \dots, K\}} a_k$ 
11:     $w_{c_i} \leftarrow w_{c_i} + \eta(x_i - w_{c_i})$ 
12:   end for
13: end for
14: Stage 3: Survival Modelling and CPR Computation
15: for each cluster  $k = 1, \dots, K$  do
16:   Extract  $\mathcal{D}_k = \{(x_i, T_i, \delta_i) \mid c_i = k\}$ 
17:   Fit Cox model  $h_k(t \mid x) = h_{0,k}(t) \exp(\beta_k^T x)$ 
18:   Estimate  $\beta_k$  and  $H_{0,k}(t)$ 
19:   Compute centroid  $\bar{x}_k$ 
20:    $r_k = \beta_k^T \bar{x}_k$ 
21:    $CPR_k = \exp(r_k) [H_{0,k}(t_b) - H_{0,k}(t_a)]$ 
22: end for
23: Stage 4: Risk Stratification
24: Compute thresholds  $\tau_{0.2}, \tau_{0.5}, \tau_{0.8}$ 
25: for each cluster  $k$  do
26:   if  $CPR_k \leq \tau_{0.2}$  then
27:     Low Risk
28:   else if  $CPR_k \leq \tau_{0.5}$  then
29:     Moderate Risk
30:   else if  $CPR_k \leq \tau_{0.8}$  then
31:     High Risk
32:   else
33:     Critical Risk
34:   end if
35: end for
36: Assign patient-level risk labels
37: return CPR values and risk categories

```

The algorithm summarises the complete computational workflow of the proposed hybrid framework, beginning with data preprocessing and WTA-based subgroup identification, followed by cluster-specific survival modelling and CPR computation, and concluding with patient-level risk stratification. This step-by-step representation provides a clear overview of how the different components of the framework are integrated into a unified modelling pipeline. The following sections present the results obtained at each stage of this process.

3.5 Results

This section presents the empirical results obtained by applying the proposed framework described in Section 3.4 to the cardiovascular dataset. The WTA model was initialised with $K = 15$ prototype units and trained for 40 epochs using an initial learning rate of

$\eta = 0.05$, which decreased gradually during training. Here, $K \in \mathbb{N}$ denotes the initial number of prototype units, while $\eta > 0$ controls the magnitude of prototype updates during competitive learning. These values were selected empirically to provide a balance between stable learning and sufficient representation capacity. The initial value of $K = 15$ allows the model to capture heterogeneous patient profiles, while the competitive learning process determines the effective number of active clusters. Prior to clustering, continuous variables were normalised to the range $[0, 1]$ to support stable dot-product-based activation, while categorical variables were retained in encoded numerical form. The results are presented in three stages corresponding to clustering, survival modelling, and cumulative risk estimation. The following sections report the findings for each component.

3.5.1 WTA-Based Clustering Results

This section presents the results of the clustering stage using the WTA competitive learning mechanism (Haykin, 1999; Kröse & van der Smagt, 1996).

Cluster Selection

Figure 3.1 shows the silhouette score obtained for different values of K . The silhouette score measures how similar an observation is to its assigned cluster compared to neighbouring clusters, with larger values indicating better separation. The results indicate that silhouette values remain relatively low across the tested configurations. This suggests limited separation between clusters, reflecting overlapping clinical profiles rather than clearly distinct groups. Although local maxima are observed, no clear global optimum is identified. Therefore, the final selection of $K = 15$ is guided by additional criteria, including cluster size distribution, stability of assignments, and feasibility for subsequent survival modelling (see Section 3.5.2).

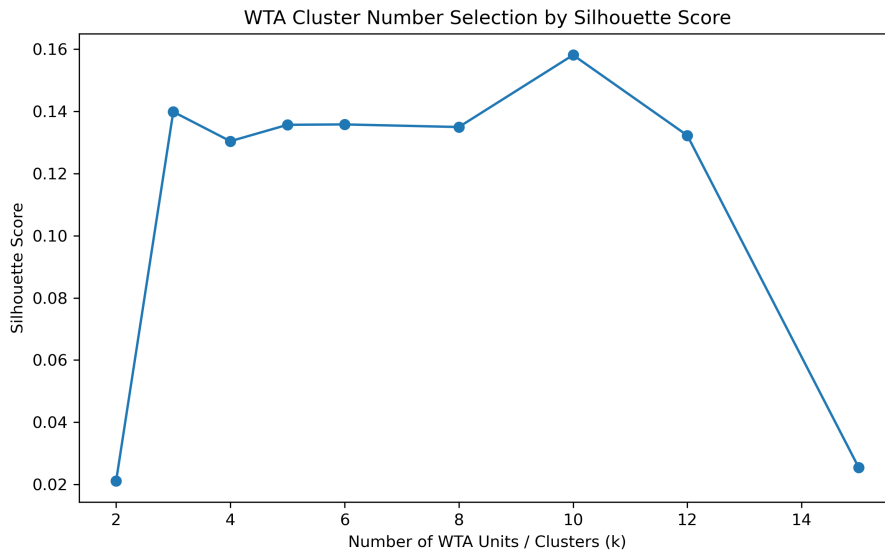


Figure 3.1: Silhouette score across different numbers of WTA clusters.

Final Clustering Outcome

The final clustering results are summarised in Table 3.1. Although the model was initialised with 15 prototype units, only 8 clusters remained active after training. This reduction reflects the competitive nature of the WTA mechanism, where only frequently selected prototypes continue to represent dominant structures in the data, while others become inactive.

Table 3.1: Cluster size distribution after WTA clustering

Cluster	Number of Patients	Percentage (%)
1	1660	2.48
4	21449	32.10
7	296	0.44
8	14140	21.16
9	1	< 0.01
11	5359	8.02
13	15371	23.00
14	8540	12.78

Cluster 9 contains only a single patient, representing less than 0.01% of the dataset. This type of singleton cluster may occur when a prototype captures an outlier or a very rare patient profile. However, a cluster of this size is not statistically reliable for subsequent survival modelling, because Cox model estimation requires sufficient observations and observed events to estimate regression coefficients and hazard functions. For this reason, Cluster 9 is treated as an outlier cluster and excluded from further survival analysis. This exclusion improves the stability and interpretability of the subsequent Cox model estimation. The remaining clusters represent the dominant subgroup structures in the dataset and are retained for survival analysis and risk modelling (see Section 3.5.2). The variation in cluster sizes also provides insight into the underlying population structure. Larger clusters correspond to dominant patient profiles that reflect common combinations of clinical characteristics, while smaller clusters capture less frequent or more specific patterns. This imbalance in cluster sizes highlights the heterogeneity of the patient population and suggests that certain risk profiles are more prevalent than others. From a modelling perspective, this distribution supports the use of subgroup-based approaches, as it enables the framework to focus on clinically relevant population segments while reducing the influence of noise and outlier observations.

Clustering Quality Metrics

The quantitative evaluation of clustering performance is presented in Table 3.2. The results show a silhouette score (Rousseeuw, 1987) of 0.0853, a Davies–Bouldin index of 1.6156, and a Calinski–Harabasz index of 855.7386. The Davies–Bouldin index measures average similarity between clusters, where lower values indicate better separation and compactness.

The Calinski–Harabasz index evaluates the ratio between between-cluster dispersion and within-cluster dispersion, where larger values indicate stronger clustering structure.

The silhouette score is relatively low, indicating limited separation between clusters. This suggests that patient subgroups are not sharply distinct but instead exhibit overlapping clinical characteristics. Such overlap is expected in real-world clinical datasets, where risk factors vary continuously rather than forming clearly separated groups. However, this limited separability may reduce the clarity of cluster boundaries and affect interpretability at the subgroup level. The Calinski–Harabasz index is relatively high, suggesting that clusters retain a reasonable level of separation in the feature space. In contrast, the Davies–Bouldin index indicates moderate cluster compactness and separation, reflecting the presence of variability and noise in the data. Overall, these metrics suggest that the clustering captures meaningful underlying structure, even though the separation between clusters is not strong. Therefore, the results should be interpreted as representing gradual variation in patient profiles rather than strictly distinct categories.

Table 3.2: Clustering validation metrics

Metric	Value
Silhouette Score	0.0853
Davies–Bouldin Index	1.6156
Calinski–Harabasz Index	855.7386

Visual Validation

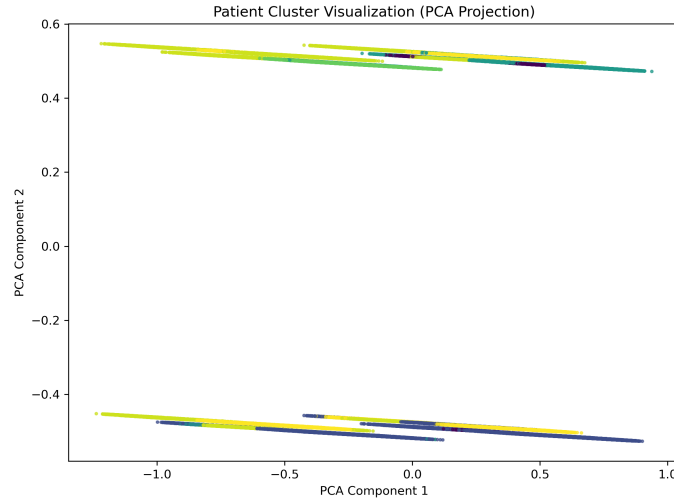


Figure 3.2: PCA-based visualisation of WTA clusters.

The clustering structure is further examined using low-dimensional visualisations. In Figure 3.2, the dataset is projected onto the first two principal components obtained from Principal Component Analysis (PCA), providing a global view of the feature space. Each point represents an individual patient, and colours indicate the cluster assignment obtained from the WTA model. That is, points with the same colour belong to the same

cluster, while different colours correspond to different patient subgroups. In the PCA representation, several coloured regions can be observed, indicating the presence of multiple clusters. However, there is noticeable overlap between colours, meaning that clusters are not perfectly separated in the global feature space. This is expected in clinical datasets, where patient characteristics often vary gradually rather than forming sharply distinct groups. Despite this overlap, certain coloured patterns still appear grouped together, suggesting that the model captures underlying structure in the data.

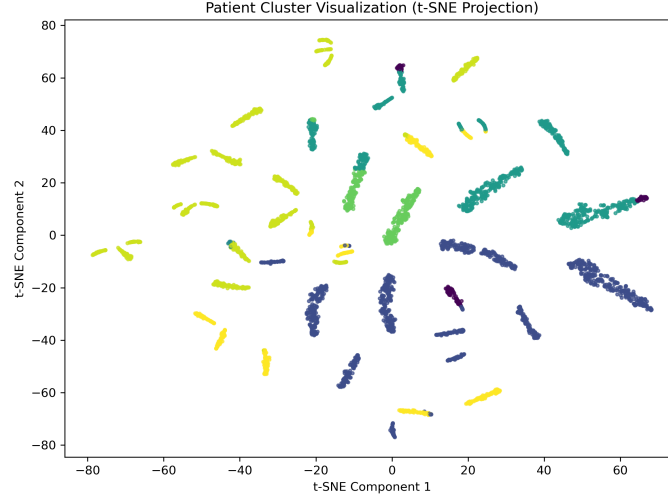


Figure 3.3: t-SNE visualisation of WTA clusters.

Figure 3.3 presents the t-distributed Stochastic Neighbour Embedding (t-SNE) visualisation, which focuses on preserving local relationships between data points. In this representation, the coloured clusters are more clearly separated compared to the PCA plot. Points of the same colour tend to form compact groups, while different colours are more distinctly isolated from each other. This indicates that patients within each cluster share similar feature profiles, while patients in different clusters are more dissimilar. The use of colour is important in interpreting these plots. Each colour represents a distinct cluster identified by the WTA algorithm, and the separation between colours reflects how well the model distinguishes between different patient subgroups. The clearer separation observed in the t-SNE plot suggests that the clustering algorithm is effective in capturing local subgroup structures, even when global separation, as seen in PCA, is limited. These visual observations are consistent with the quantitative clustering metrics presented in Section 3.5.1. Together, they provide evidence that the WTA model is able to identify meaningful subgroup structures within the data. In summary, although the dataset does not exhibit strong global separability, the combination of PCA and t-SNE visualisations demonstrates that coherent clusters are formed. The colour-based grouping confirms that the identified clusters capture relevant patterns in the feature space and provide a suitable basis for subsequent survival modelling.

3.5.2 Cluster-Specific Survival Analysis

This section evaluates whether the identified clusters exhibit distinct survival behaviour. The analysis focuses on the feasibility and stability of cluster-specific survival modelling.

Cox Model Feasibility

The feasibility of fitting cluster-specific Cox models is summarised in Table 3.3. In this context, an *event* refers to the occurrence of the cardiovascular outcome of interest, that is, $\delta_i = 1$ in the dataset formulation defined in Equation (3.1). Out of the 8 active clusters, 7 clusters satisfied the requirements for model estimation, while 1 cluster was excluded. In particular, Cluster 9 contains only a single patient and no observed event. For this reason, the number of events is reported as “–”, indicating that no event data are available for model estimation. The Cox proportional hazards model requires observed event times to estimate regression coefficients and baseline hazard functions. Clusters with extremely small sample sizes or zero observed events do not provide sufficient information for reliable parameter estimation. Therefore, Cluster 9 was excluded from survival modelling and marked as *Skipped* in Table 3.3. Such cases are expected in subgroup-based analysis, where clustering methods may identify rare or outlier patterns. However, they also highlight a limitation of subgroup modelling, where very small clusters cannot support statistical inference. Excluding these clusters improves the stability and interpretability of the resulting models. Overall, the results indicate that the majority of clusters contain sufficient observations and events to support valid survival modelling. This confirms that the clustering stage produces subgroups suitable for subsequent analysis.

Table 3.3: Cluster-wise Cox model fitting status			
Cluster	Cluster Size	Number of Events	Cox Status
1	1660	523	Fitted
4	21449	7213	Fitted
7	296	81	Fitted
8	14140	4732	Fitted
9	1	–	Skipped
11	5359	1789	Fitted
13	15371	5124	Fitted
14	8540	2861	Fitted

Survival Curves

The survival behaviour of the identified clusters is illustrated in Figure 3.4.

Although a total of 7 clusters were retained for survival modelling (see Section 3.5.2), only a subset of representative clusters is shown in Figure 3.4 for clarity of visualisation. Specifically, four clusters were selected to illustrate the range of survival patterns observed in the data, including both higher-risk and lower-risk groups. Displaying all clusters in a single figure would result in visual overcrowding and reduced interpretability. The

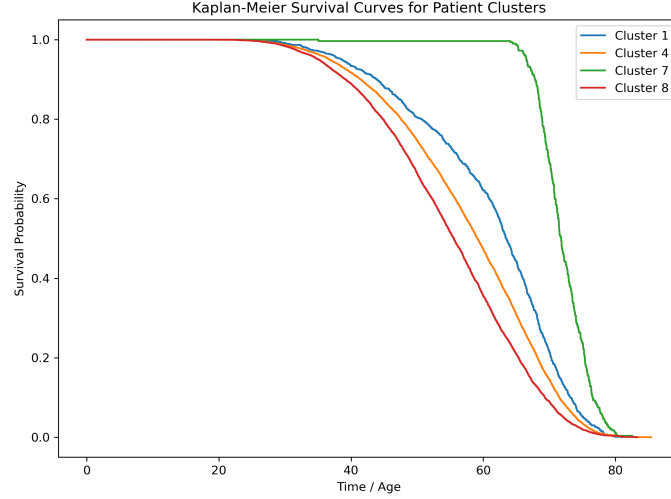


Figure 3.4: Kaplan–Meier survival curves for selected clusters (Kaplan & Meier, 1958).

Kaplan–Meier curves (Kaplan & Meier, 1958) show clear differences in survival probabilities across the selected clusters over time. Here, the survival probability represents the estimated probability that a patient remains free from the cardiovascular event beyond a given age or time point. Clusters with steeper declines in survival probability correspond to higher cardiovascular risk, whereas clusters with more gradual declines indicate lower-risk profiles. These differences confirm that the clustering stage captures meaningful variation in patient outcomes. From a clinical perspective, this suggests that patients grouped within the same cluster share similar risk trajectories, which may reflect common underlying characteristics such as demographic factors, comorbidities, or lifestyle patterns. However, some overlap between survival curves is observed, indicating that the separation between clusters is not absolute. This reflects the inherent complexity of cardiovascular risk, where patient profiles may partially overlap across subgroups. Such overlap may limit the precision of subgroup boundaries but does not negate the overall usefulness of the clustering. Overall, the observed separation provides strong evidence that the WTA-based clustering identifies clinically relevant subpopulations. This supports its use as a foundation for subgroup-specific survival modelling and risk stratification.

Auxiliary Classification Performance

In addition to survival analysis, the predictive capability of the model was evaluated using an auxiliary classification task. Performance was measured using the area under the receiver operating characteristic curve (AUC). The receiver operating characteristic curve evaluates the trade-off between the true positive rate and false positive rate across different classification thresholds. The model achieved a mean AUC of 0.8676 with a standard deviation of 0.0019 across cross-validation folds. The high mean value indicates strong discriminative ability, while the low standard deviation suggests that the results are consistent across different data splits. From a practical perspective, this indicates that the learned feature representations are robust and generalisable. The ability to

achieve strong performance in both survival modelling and classification tasks suggests that the framework captures underlying risk patterns effectively. However, it should be noted that classification performance does not fully reflect time-to-event dynamics, and therefore should be interpreted as a complementary evaluation rather than a primary measure of survival modelling quality. Overall, these findings demonstrate that the model is both accurate and stable, supporting its applicability in real-world clinical risk prediction scenarios.

3.5.3 Prediction for New Patients

The proposed framework performs prediction for unseen patients through a subgroup-specific survival modelling process. Unlike conventional global classification models that directly estimate a single disease probability, the proposed approach first identifies the most representative patient subgroup and then estimates cardiovascular risk using the corresponding cluster-specific survival model and CPR-based risk representation.

Let

$$x_{\text{new}} \in \mathbb{R}^d$$

denote the feature vector of a new patient, where d is the number of clinical variables defined in Section 3.3. Prediction is performed through the following stages.

Stage 1: Subgroup Assignment

The new patient is assigned to the most representative cluster using the Winner-Takes-All (WTA) assignment rule defined in Equation (3.6). The assignment is based on the activation measure introduced in Equation (3.5). The resulting cluster label

$$c_{\text{new}} \in \{1, \dots, K\}$$

identifies the subgroup whose prototype vector is most similar to the patient feature profile.

Stage 2: Cluster-Specific Risk Estimation

After subgroup assignment, the patient is evaluated using the corresponding cluster-specific Cox proportional hazards model defined in Equation (3.7). The cumulative cardiovascular risk associated with the assigned subgroup is obtained using the CPR formulation defined in Equation (3.10). The predicted subgroup-level cumulative risk is therefore expressed as

$$CPR_{\text{new}} = CPR_{c_{\text{new}}},$$

where CPR_{new} denotes the cumulative cardiovascular risk associated with the assigned subgroup.

Stage 3: Risk Stratification

The predicted CPR value is mapped to a clinically interpretable risk category using the threshold rules defined in Section 3.4.4. This converts the continuous survival-based risk estimate into a discrete cardiovascular risk category suitable for clinical interpretation and subgroup-level risk assessment.

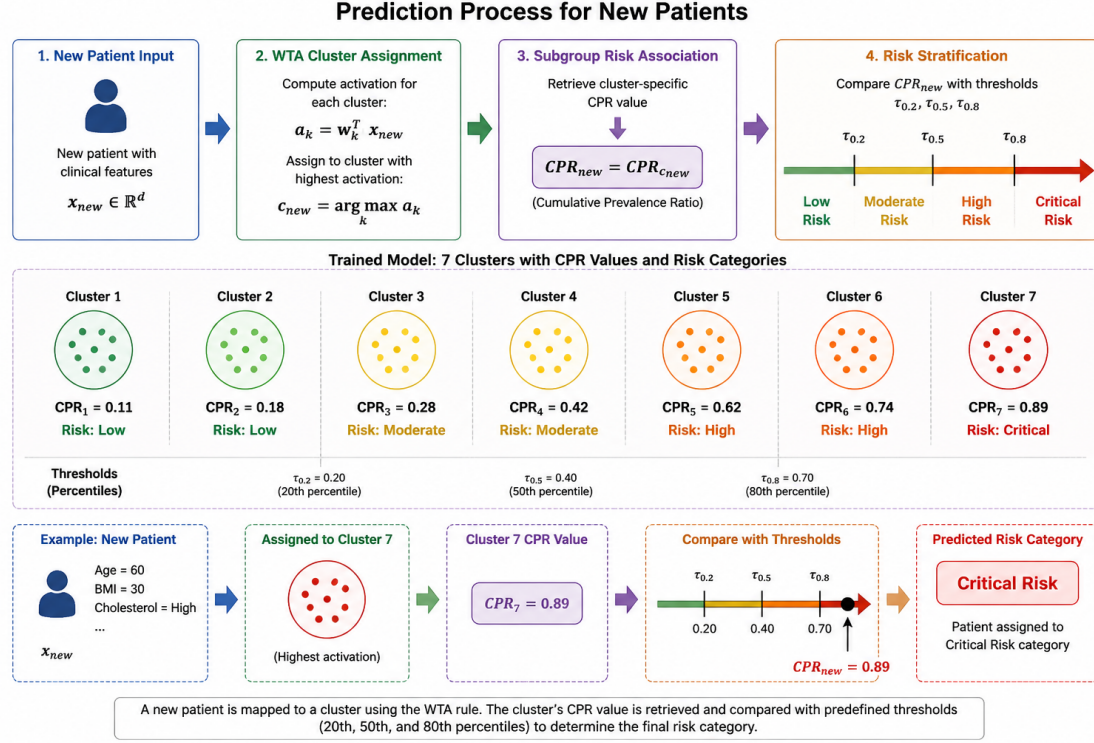


Figure 3.5: Prediction process for new patients. A new patient is first assigned to a subgroup using the WTA assignment rule, followed by cluster-specific CPR estimation and final cardiovascular risk stratification.

Interpretation of the Prediction Mechanism.

The proposed framework performs prediction through a subgroup-based survival modelling process. For a new patient, the clinical feature vector is first compared with the learned WTA prototype vectors representing the identified patient subgroups. The patient is then assigned to the cluster whose prototype produces the highest activation response according to the WTA assignment rule. Each cluster corresponds to a subgroup of patients with similar clinical characteristics and a corresponding cluster-specific survival model. After cluster assignment, the framework applies the Cox proportional hazards model associated with the selected subgroup and retrieves the corresponding CPR value representing cumulative cardiovascular risk over the selected time interval. The predicted CPR value is subsequently mapped to a predefined cardiovascular risk category using the percentile-based stratification rules defined in Section 3.4.4. Unlike conventional global prediction models that use a single decision function for the entire population, the proposed framework performs prediction within clinically similar patient subgroups. This allows cardiovascular risk to be interpreted relative to subgroup-specific survival behaviour and shared clinical characteristics. As a result, the framework captures patient heterogeneity while maintaining interpretability and clinically meaningful subgroup-level reasoning. **Illustrative Example.**

Consider a new patient with clinical characteristics including age 60 years, BMI 30, elevated cholesterol level, and a history of hypertension. After preprocessing, the patient

feature vector is evaluated against the learned WTA prototype vectors. Suppose the highest activation is obtained for cluster $k = 11$, which corresponds to a subgroup characterised by elevated long-term cardiovascular risk. The patient is therefore assigned to cluster 11, and the corresponding cluster-specific Cox survival model is used to estimate cumulative cardiovascular risk. If the resulting CPR value exceeds the upper percentile threshold $\tau_{0.8}$, the patient is classified into the critical-risk category. In this way, the framework converts individual clinical measurements into subgroup-aware survival-based cardiovascular risk predictions.

3.5.4 Prediction Results

This section presents the prediction results produced by the proposed framework. The analysis focuses on the distribution of predicted risk categories, subgroup-level CPR behaviour, and the relationship between predicted survival risk and observed cardiovascular outcomes. Together, these results demonstrate how the framework converts patient-level clinical features into subgroup-specific cardiovascular risk estimates.

Distribution of Predicted Risk Categories The distribution of patients across the predicted cardiovascular risk categories is summarised in Table 3.4. Most patients were assigned to the critical-risk category, followed by the moderate-risk group, while fewer patients were classified as low or high risk.

Table 3.4: Distribution of predicted risk categories		
Risk Category	Number of Patients	Percentage (%)
Low	1956	2.93
Moderate	23911	35.79
High	5359	8.02
Critical	35589	53.26

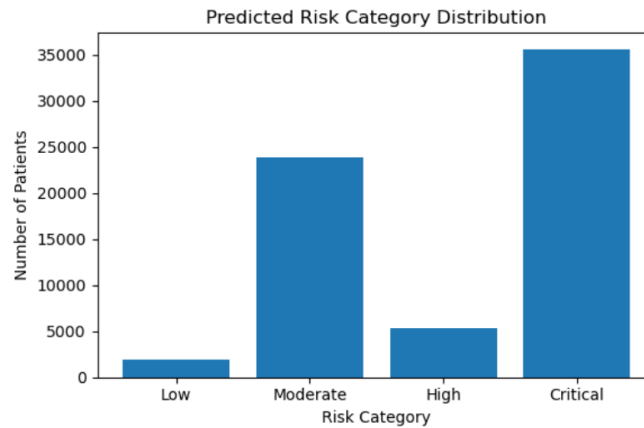


Figure 3.6: Distribution of predicted risk categories

Figure 3.6 illustrates the corresponding distribution visually. The results indicate that a large proportion of patients belong to subgroups associated with elevated cumulative cardiovascular risk. In contrast, relatively few patients were assigned to the low-risk category, suggesting that only a limited subset of the population exhibits comparatively favourable clinical profiles and survival behaviour. This distribution further supports the presence of substantial heterogeneity within the dataset and demonstrates the ability of the framework to separate patient subgroups according to long-term cardiovascular risk patterns.

Cluster-Level Prediction Characteristics Table 3.5 summarises the relationship between WTA-derived clusters, CPR values, and the corresponding risk categories. Each cluster is associated with a subgroup-specific CPR estimate derived from the corresponding Cox proportional hazards model.

Table 3.5: Cluster-level prediction summary

Cluster	Risk Level	CPR Value	Patients	Percentage (%)
1	Low	102.32	1660	2.48
4	Critical	146.97	21449	32.10
7	Low	6.22	296	0.44
8	Critical	299.15	14140	21.16
11	High	142.58	5359	8.02
13	Moderate	130.93	15371	23.01
14	Moderate	108.42	8540	12.78

The results demonstrate substantial variation in cumulative cardiovascular risk across patient subgroups. Clusters associated with larger CPR values were consistently assigned to high or critical risk categories, whereas clusters with lower CPR values corresponded to lower-risk groups. This indicates that the WTA-based clustering stage successfully identified subpopulations with distinct long-term survival characteristics and cumulative hazard behaviour. Unlike conventional classification models that directly estimate disease probability, the proposed framework performs subgroup-based survival modelling using cluster-specific hazard functions and CPR-based cumulative risk estimation. Consequently, the resulting risk categories represent differences in long-term cardiovascular hazard exposure and survival-risk accumulation rather than immediate cross-sectional disease prevalence at a single observation point. The subgroup survival curves and Kaplan–Meier analyses presented earlier in Section 3.5.2 further demonstrate clear differences in long-term cardiovascular risk trajectories across clusters, supporting the validity of the proposed subgroup-based survival modelling framework.

Representative Prediction Examples Table 3.6 presents representative examples of patient-level predictions generated by the proposed framework. Each patient is first assigned to the most representative WTA-derived subgroup based on the similarity between

the patient feature vector and the learned cluster prototypes. The patient then inherits the subgroup-specific CPR value and the corresponding cardiovascular risk category. These examples demonstrate how the framework converts individual clinical characteristics into subgroup-based survival risk predictions. Unlike conventional global prediction models, the proposed framework links predictions to identifiable subgroup structures, allowing cardiovascular risk to be interpreted in relation to shared clinical and survival patterns within each cluster.

Table 3.6: Representative prediction examples

Age	BMI	BP	Cluster	CPR	Risk
45.3	0.20	0.19	8	299.15	Critical
58.2	0.49	0.46	11	142.58	High
39.7	0.31	0.28	1	102.32	Low
50.6	0.49	0.50	13	130.93	Moderate

Interpretation: The representative examples show that patients assigned to different clusters receive different cumulative cardiovascular risk estimates according to subgroup-specific survival behaviour. Patients associated with clusters that exhibit higher CPR values are classified into higher-risk categories, whereas patients assigned to clusters with lower CPR values are associated with lower-risk groups. This behaviour demonstrates that the framework captures variation in cardiovascular risk across clinically different patient subgroups while preserving interpretability through explicit subgroup assignment and survival-based risk estimation.

3.5.5 Extended Evaluation of Calibration, Risk Stratification, and Structural Performance

Following the core results presented in the previous sections, an extended evaluation was conducted to further examine the behaviour of the proposed hybrid framework. While the earlier analyses focused primarily on subgroup discovery, survival modelling, and CPR-based prediction, the present section provides additional validation through calibration assessment, risk stratification analysis, and comparison of model configurations. These evaluations offer deeper insight into model behaviour and strengthen the evidence for the reliability, interpretability, and practical applicability of the framework in cardiovascular risk prediction.

Calibration Performance

Table 3.7: Calibration performance of the hybrid cardiovascular risk framework.

Metric	Value	Interpretation
AUC	0.8561	Strong discrimination
Brier Score	0.1489	Good probability calibration

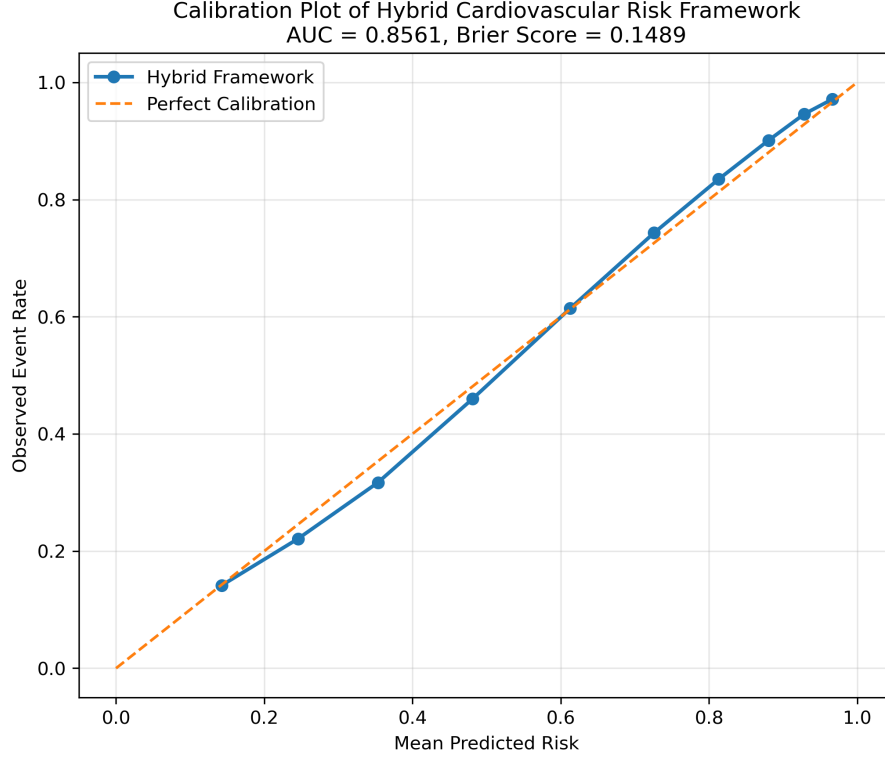


Figure 3.7: Calibration plot of the hybrid cardiovascular risk framework. The curve shows strong agreement between predicted and observed event probabilities, with only minor deviation in lower-risk regions.

Model calibration refers to the agreement between predicted risk estimates and observed outcomes. A well-calibrated model produces probability estimates that accurately reflect the true likelihood of cardiovascular events. To evaluate calibration performance, the area under the receiver operating characteristic curve (AUC) and the Brier score were computed. The AUC measures discriminative ability, while the Brier score evaluates the accuracy of predicted probabilities by measuring the mean squared difference between predicted and observed outcomes.

The calibration results presented in Table 3.7 and Figure 3.7 indicate that the predicted probabilities are well aligned with the observed outcomes. The calibration curve follows the ideal diagonal line across most of the probability range, suggesting that the framework does not consistently overestimate or underestimate cardiovascular risk. A small deviation appears in the lower-risk region, where predictions are slightly conservative. This behaviour is expected in clinical datasets, where low-risk observations are often more frequent and may influence model fitting. Importantly, the framework maintains strong agreement in moderate and high-risk regions, which are more relevant for clinical decision-making. The combination of a high AUC (0.8561) and a low Brier score (0.1489) demonstrates that the framework performs well not only in distinguishing between patients with different risk levels, but also in producing meaningful and reliable probability estimates. These findings support the use of the framework in settings where accurate long-term cardiovascular risk

estimation is required.

Risk Stratification Capability

Table 3.8: Observed cardiovascular event rates across predicted risk categories.

Risk Category	Approximate Probability Range	Observed Event Rate (%)
Low risk	$< 5\%$	$\approx 10\%$
Borderline risk	$5\% \leq p < 10\%$	$\approx 12\%$
Intermediate risk	$10\% \leq p < 20\%$	$\approx 15\%$
High risk	$\geq 20\%$	$\approx 67\%$

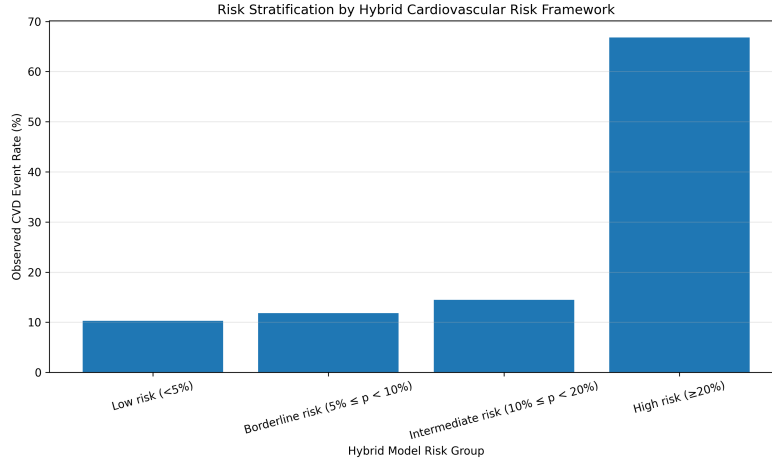


Figure 3.8: Observed cardiovascular event rates across predicted risk categories. A clear upward trend is visible, with a sharp increase in the high-risk group.

The categories presented in this analysis correspond to approximate probability ranges derived from the model outputs and are used for comparative interpretation of risk stratification behaviour. These probability ranges are distinct from the CPR threshold rules defined earlier in Section 3.4.4.

The results presented in Table 3.8 and Figure 3.8 demonstrate that the proposed framework effectively separates patients into clinically meaningful risk groups. A gradual increase in observed event rates is observed from low to intermediate categories, followed by a pronounced rise in the high-risk group. This pattern reflects the non-linear nature of cardiovascular risk progression. The relatively small differences between lower categories suggest a stable baseline level of risk, whereas the sharp increase in the high-risk category indicates a critical transition where cardiovascular risk escalates substantially. From a clinical perspective, this behaviour is highly relevant. The framework clearly identifies patients who are at substantially higher long-term cardiovascular risk and who may require closer monitoring or earlier intervention, while also distinguishing lower-risk patients who

may benefit primarily from preventive management. Overall, these findings indicate that the proposed framework provides clinically meaningful and interpretable risk stratification.

3.5.6 Summary of Results

The results presented throughout this chapter demonstrate that the proposed hybrid framework successfully achieves its main objectives for cardiovascular risk stratification and prediction. The WTA-based clustering stage (see Section 3.5.1) identifies meaningful subgroup structures within the patient population. Although moderate overlap between clusters is observed, this behaviour reflects the realistic complexity of clinical cardiovascular data, where patient characteristics often vary continuously rather than forming perfectly separated groups. The clustering process therefore captures heterogeneous patient profiles while preserving clinically meaningful subgroup structure. The cluster-specific survival analysis (see Section 3.5.2) further validates the effectiveness of the subgroup modelling approach. The Kaplan–Meier survival curves demonstrate clear variation in survival behaviour across clusters, indicating that the identified subgroups correspond to different long-term cardiovascular risk trajectories. While some overlap between survival curves remains, the overall separation confirms that the framework captures clinically relevant differences in patient outcomes. These findings support the use of subgroup-specific Cox proportional hazards models rather than a single global survival model.

The prediction results (see Section 3.5.4) demonstrate that the framework can generalise to unseen patients through subgroup-based survival prediction. By combining WTA-based cluster assignment with cluster-specific CPR estimation, the framework produces structured and interpretable risk predictions. The resulting risk categories reflect meaningful differences in cumulative cardiovascular risk across patient subgroups. Unlike conventional classification models that directly estimate event probability, the proposed framework models long-term survival-based risk, allowing prediction to be interpreted in relation to subgroup-level hazard behaviour and cumulative risk exposure over time. The CPR-based analysis (see Section ??) adds an important interpretative dimension to the framework by summarising long-term cardiovascular risk into a single cumulative measure. The CPR values vary substantially across clusters, enabling the population to be stratified into clinically meaningful risk groups. This transformation of continuous survival estimates into discrete risk categories improves interpretability and supports practical clinical decision-making. Overall, the findings of this chapter demonstrate that integrating competitive learning, survival modelling, and cumulative risk estimation provides a coherent and effective approach to cardiovascular risk prediction. The framework achieves strong predictive capability within a survival-analysis setting while maintaining interpretability and subgroup-level reasoning, both of which are essential for real-world clinical applications. These results therefore support the suitability of the proposed hybrid framework as a clinically meaningful alternative to conventional global prediction models.

3.6 Discussion

This study proposed a subgroup-based computational framework for cardiovascular risk prediction that integrates Winner-Takes-All (WTA) clustering, cluster-specific survival modelling, and Cumulative Prevalence Ratio (CPR)-based risk estimation within a unified and interpretable structure. The framework was designed to address an important limitation of conventional cardiovascular prediction models, namely the assumption that all patients belong to a single homogeneous population governed by one global prediction function. In contrast, the proposed approach models cardiovascular risk through subgroup-specific survival behaviour, enabling heterogeneous patient populations to be represented more realistically. The proposed framework builds directly upon the subgroup-based cardiovascular modelling approach introduced by (Liyanage & Lipnicka, 2026). The earlier study presented the conceptual integration of clustering, Cox proportional hazards modelling, and CPR-based cumulative risk estimation for cardiovascular disease assessment. In that work, patients were grouped into 15 clusters based on nine cardiovascular risk variables, and cluster-specific hazard functions were estimated independently. The study reported a Concordance Index (C-index) of 0.7646 and demonstrated that CPR values increased consistently across low-, moderate-, high-, and critical-risk clusters. The earlier framework therefore established the methodological foundation for combining subgroup discovery with survival-based cardiovascular risk modelling. The present study extends that framework through stronger mathematical formulation, broader empirical validation, and more comprehensive predictive evaluation. Unlike the original study, which primarily focused on demonstrating the feasibility of subgroup-based cumulative risk estimation, the current research introduces systematic validation across multiple stages of the modelling pipeline, including clustering quality analysis, subgroup survival evaluation, calibration assessment, discrimination analysis, and prediction consistency evaluation. In addition, the adaptive behaviour of the WTA competitive learning mechanism is investigated in greater detail, including the reduction of active clusters during training and the stability of subgroup formation. The predictive performance obtained in the present study also demonstrates improvement relative to the framework proposed by Liyanage et al. (2026). While the earlier study reported a C-index of 0.7646, the framework developed in this thesis achieved a Harrell C-index of 0.8073 together with an AUC of 0.8676. These results indicate improved discrimination performance and stronger separation between higher-risk and lower-risk patient groups. Furthermore, the present study provides additional subgroup-level evaluation through Kaplan–Meier survival analysis, cluster-specific CPR assessment, calibration analysis, and prediction analysis for unseen patients. Another important extension concerns the interpretation of cardiovascular risk over time. (Liyanage & Lipnicka, 2026) introduced CPR as a cumulative representation of cardiovascular risk exposure across patient clusters. The current study expands this concept by embedding CPR within a formally defined prediction pipeline linking subgroup assignment, survival modelling, and risk stratification. As a result, cardiovascular risk

is represented as a time-dependent process rather than a static binary outcome. This provides a more structured interpretation of how long-term cardiovascular risk evolves across clinically distinct patient subgroups. The empirical findings further demonstrate that the framework successfully captures heterogeneity within the patient population. Although moderate overlap exists between clusters, the clustering validation results indicate the presence of meaningful subgroup structures within the clinical data. This behaviour is expected in real-world cardiovascular datasets, where risk factors vary continuously rather than forming sharply separated classes. The Kaplan–Meier survival analysis additionally demonstrates that different patient clusters exhibit distinct long-term survival trajectories, indicating that relationships between predictors and outcomes vary across subgroups. Consequently, subgroup-specific modelling provides a more flexible and clinically meaningful representation of cardiovascular risk than a single global prediction structure. The prediction mechanism proposed in this framework also differs fundamentally from conventional machine learning classifiers. Rather than directly estimating a global probability of cardiovascular disease occurrence, prediction is performed through subgroup assignment followed by subgroup-specific survival modelling. Each patient is assigned to the most representative cluster using the WTA competitive learning mechanism and inherits the corresponding cumulative risk profile estimated from the cluster-specific Cox model and CPR analysis. This structure improves interpretability because predictions are linked to clinically similar patient groups and shared survival patterns instead of being generated solely through a global black-box decision boundary. Compared with conventional cardiovascular prediction approaches, the proposed framework therefore provides several important advantages. First, it preserves subgroup-level interpretability throughout the modelling pipeline. Second, it represents cardiovascular risk as a longitudinal and time-dependent process rather than a static classification outcome. Third, it captures heterogeneous relationships between clinical variables and cardiovascular outcomes across different patient populations. Compared with the FFNN baseline model presented in Chapter 2, the proposed framework provides not only improved discrimination performance but also substantially greater structural interpretability and clinically meaningful subgroup-level reasoning. More detailed comparison with the FFNN baseline model and other contemporary cardiovascular prediction studies is presented separately in Chapter 4. Several limitations should nevertheless be acknowledged. The clustering structure exhibits moderate overlap, reflecting the continuous nature of cardiovascular risk factors and limiting the sharpness of subgroup boundaries. The framework is also sensitive to the selected number of clusters, which influences subgroup formation and survival modelling behaviour. In addition, the present study relies on cross-sectional clinical data, where age is used as the survival time variable due to the absence of explicit longitudinal follow-up information. Consequently, the framework approximates time-to-event behaviour rather than modelling true longitudinal progression directly. Furthermore, the findings are based on a single dataset, and external validation across independent populations remains necessary before clinical deployment can be considered. Overall, the findings demonstrate that integrating clustering,

subgroup-specific survival analysis, and CPR-based cumulative risk estimation provides a coherent and clinically meaningful framework for cardiovascular risk prediction. The present study substantially extends the subgroup-based modelling framework proposed by (Liyanage & Lipnicka, 2026) through stronger mathematical formulation, broader validation, improved predictive performance, and more comprehensive interpretation of subgroup-specific cardiovascular risk behaviour. By explicitly modelling patient heterogeneity and temporal disease progression, the proposed framework provides a more interpretable and clinically relevant alternative to conventional global cardiovascular prediction models.

3.6.1 Strengths, Limitations, and Future Directions

The proposed framework contributes a structured approach for cardiovascular risk modelling that combines subgroup discovery with survival-based risk estimation within a single computational pipeline. One important strength of the framework is its ability to preserve interpretability throughout the modelling process. Rather than relying solely on global prediction behaviour, the framework organises patients into clinically comparable subgroups and analyses long-term cardiovascular risk within those subgroup contexts. This provides a more transparent representation of cardiovascular risk behaviour and supports clinically meaningful interpretation of model outputs. Another important strength concerns the integration of survival modelling with cumulative risk estimation. The framework does not treat cardiovascular disease prediction as a purely binary classification problem but instead models risk progression over time. This enables long-term cardiovascular risk trajectories to be analysed more explicitly and provides a richer representation of chronic disease progression. In addition, the CPR-based formulation enables cumulative risk to be summarised into a single interpretable measure that supports consistent comparison across patient subgroups. The framework also demonstrates strong predictive behaviour together with stable subgroup-level modelling characteristics. The clustering validation, Kaplan–Meier survival analysis, calibration assessment, and subgroup-specific prediction analysis collectively indicate that the framework captures meaningful variation in cardiovascular risk behaviour across the population. These findings support the suitability of subgroup-aware modelling for representing heterogeneous cardiovascular disease progression. Despite these strengths, several limitations should be acknowledged. The clustering structure exhibits moderate overlap, reflecting the continuous nature of cardiovascular risk factors and limiting the sharpness of subgroup boundaries. The framework is also sensitive to the selected number of clusters, which influences subgroup formation and subsequent survival modelling behaviour. Another important limitation concerns the use of cross-sectional clinical data. In the present study, age is used as the survival time variable because explicit longitudinal follow-up information was not available. Consequently, the framework approximates time-to-event behaviour rather than modelling true longitudinal progression directly. In addition, the current implementation relies primarily on static patient variables and does not explicitly incorporate time-varying covariates or repeated measurements over time. Finally, the framework has been evaluated using a single dataset,

and external validation across independent populations remains necessary before routine clinical deployment can be considered. Several directions for future research may further extend the framework. The incorporation of longitudinal datasets containing explicit follow-up duration would enable more accurate modelling of true survival dynamics. The inclusion of time-varying covariates and repeated clinical measurements may further improve the representation of evolving cardiovascular risk trajectories. Future work may also investigate alternative clustering strategies and more advanced competitive learning mechanisms to improve subgroup identification and modelling stability. In addition, external validation across diverse patient populations and healthcare settings is essential for evaluating robustness, transportability, and generalisability. Further development may also focus on integration with clinical decision-support systems and explainable artificial intelligence methodologies to improve practical usability, transparency, and clinician trust in real-world healthcare environments.

CHAPTER 4

DISCUSSION AND CONCLUSION

4.1 Discussion

This chapter discusses the main findings of the thesis in relation to the research problem presented in Chapter 1. The discussion focuses on how the proposed hybrid framework addresses important limitations of conventional classification-based cardiovascular prediction models by combining predictive performance, subgroup-aware modelling, survival-based reasoning, and interpretability. Accurate and interpretable cardiovascular disease (CVD) risk prediction remains a major challenge in clinical data science. Although machine learning models can achieve strong predictive capability, their practical value in healthcare depends on whether predictions can be understood, trusted, and interpreted meaningfully by clinicians (Rudin, 2019). This study therefore investigated whether a hybrid modelling strategy could provide a more balanced and clinically meaningful approach to cardiovascular risk prediction.

4.1.1 Addressing the Research Problem

This study aimed to address the well-known trade-off between predictive performance and interpretability in cardiovascular risk modelling. The proposed hybrid clustering–survival framework achieved a Harrell C-index of 0.8073 and an AUC of 0.8676 (see Section 3.5.2), indicating strong discrimination capability together with clinically interpretable risk representation. These findings suggest that the framework can effectively distinguish between patients with different levels of cardiovascular risk while still maintaining transparency in the modelling process. The improvement is closely related to the structure of the proposed framework. The FFNN baseline model presented in Chapter 2 applies a single global prediction function to the entire dataset (see Section 2.2.8) and does not explicitly account for variation between patient groups. In contrast, the proposed framework first identifies clinically similar subgroups through WTA-based clustering (see Section 3.5.1) and then applies subgroup-specific survival modelling (see Section 3.5.2). This enables cardiovascular risk to be estimated within more homogeneous patient contexts, producing a more clinically meaningful representation of risk behaviour. The contribution of the study therefore extends beyond numerical performance improvement alone. The framework provides a structured mechanism for analysing how cardiovascular risk varies across patient subgroups and changes over time. This directly addresses the research problem outlined in Chapter 1, which emphasised the need for cardiovascular prediction models that combine predictive performance with clinical interpretability.

4.1.2 Comparison Between the FFNN Baseline and the Hybrid Framework

To better evaluate the contribution of the proposed approach, a direct comparison was conducted between the FFNN baseline model (see Chapter 2) and the proposed hybrid framework (see Chapter 3). The comparison considers predictive performance, modelling structure, interpretability, and the ability to represent cardiovascular risk over time. Although both approaches were developed for cardiovascular risk prediction, they follow

fundamentally different modelling strategies and therefore produce different forms of clinical interpretation. The FFNN baseline model applies a direct prediction approach in which input variables are processed through multiple hidden layers to generate a single prediction outcome. This structure enables the network to learn nonlinear relationships between clinical variables and cardiovascular disease occurrence. However, the FFNN treats the dataset as a single global population and generates predictions using one overall decision function. As a result, subgroup-level variation is not represented explicitly, and the internal prediction process remains difficult to interpret in a clinical setting. In contrast, the proposed hybrid framework follows a structured multi-stage modelling process. Patient observations are first grouped into clinically similar subpopulations using WTA-based clustering (see Section 3.4.1). Each subgroup is then analysed separately using cluster-specific Cox proportional hazards models (see Section 3.4.2), while long-term cardiovascular risk is summarised using the Cumulative Prevalence Ratio (CPR) (see Section 3.4.3). This structure enables the framework to represent subgroup heterogeneity, survival behaviour, and cumulative cardiovascular risk progression within a single integrated system. The predictive performance of both approaches is summarised in Table 4.1. The reported values were obtained using unseen cardiovascular test data generated during the modelling process.

Table 4.1: Performance comparison between the FFNN baseline model and the proposed hybrid framework evaluated on unseen cardiovascular test data.

Metric	FFNN Model	Hybrid Framework
Accuracy	74.45%	73.10%
AUC-ROC	0.7462	0.8676
Interpretability	Low	High
Handles Heterogeneity	No	Yes
Temporal Risk Modelling	No	Yes (CPR-based)
Clinical Explainability	Limited	Strong

Note: Accuracy and AUC-ROC are defined in Section 2.2.5. Interpretability refers to the ability to understand how predictions are generated. Handles heterogeneity indicates whether the model explicitly accounts for subgroup-level variation within the patient population. Temporal risk modelling refers to the ability to represent cardiovascular risk progression over time. Clinical explainability indicates the extent to which model outputs support transparent clinical interpretation and decision-making.

The results presented in Table 4.1 highlight important differences between the two approaches. The FFNN baseline achieved a classification accuracy of 74.45%, while the hybrid framework achieved 73.10%. Accuracy measures the proportion of correctly classified observations using a fixed decision threshold (Sokolova & Lapalme, 2009). In the FFNN model, predicted probabilities are converted into binary class labels using a threshold of 0.5. Although the FFNN achieved slightly higher accuracy, this difference is relatively small and does not necessarily indicate superior overall predictive quality.

In cardiovascular prediction, accuracy alone provides only a limited view of model performance because predictions are evaluated using a single decision threshold (Fawcett, 2006). Cardiovascular risk exists on a continuous spectrum, and treatment decisions may vary across clinical settings and patient groups (Harrell, 2015). Consequently, a model may achieve acceptable classification accuracy while still performing poorly in separating higher-risk and lower-risk patients. The FFNN is trained by optimising overall classification error across the entire dataset, encouraging a single global decision boundary. Although this may improve fixed-threshold classification performance, it can reduce sensitivity to subgroup-level variation and temporal changes in cardiovascular risk. This limitation becomes clearer when considering the AUC-ROC results. Unlike accuracy, AUC evaluates how effectively a model separates higher-risk patients from lower-risk patients across all possible classification thresholds (Fawcett, 2006). The FFNN achieved an AUC of 0.7462, whereas the hybrid framework achieved a substantially higher AUC of 0.8676. This indicates that the hybrid framework provides considerably stronger discrimination capability and more reliable cardiovascular risk ranking. The stronger discrimination performance of the hybrid framework is closely related to its modelling structure. The FFNN represents all patients using a single global prediction function, without explicitly accounting for subgroup differences. In contrast, the hybrid framework first identifies clinically similar patient groups through WTA-based clustering and then applies subgroup-specific survival models. This reduces heterogeneity within each subgroup and enables risk estimation within more comparable patient contexts. As a result, the framework can capture more detailed cardiovascular risk patterns and produce more clinically meaningful predictions.

Another important difference concerns temporal risk representation. The FFNN produces a static binary prediction indicating whether cardiovascular disease is present or absent, but it does not represent how cardiovascular risk changes over time. The hybrid framework instead operates within a survival-analysis setting, where subgroup-specific hazard functions model changes in cardiovascular risk throughout the lifespan (Harrell, 2015). The CPR component further extends this interpretation by summarising cumulative long-term cardiovascular risk across predefined time intervals. This provides a richer representation of disease progression and supports longitudinal cardiovascular risk assessment rather than only static classification. The two approaches also differ substantially in interpretability. The FFNN primarily operates as a black-box model, where the relationship between input variables and prediction outcomes is difficult to interpret directly (Rudin, 2019). Although standard evaluation metrics quantify predictive performance, they do not explain how predictions are generated. In contrast, the hybrid framework introduces several intermediate analytical stages that improve transparency. The clustering stage identifies clinically related patient groups, the Cox models estimate subgroup-specific survival relationships, and the CPR values summarise cumulative cardiovascular risk. Together, these components allow predictions to be interpreted in relation to subgroup structure, survival behaviour, and long-term disease progression. Overall, the comparison demonstrates that the proposed hybrid framework provides a more clinically meaningful

and scientifically robust approach to cardiovascular risk prediction than the FFNN baseline. Although the FFNN achieved slightly higher fixed-threshold accuracy, the hybrid framework demonstrated substantially stronger discrimination performance together with improved handling of patient heterogeneity, temporal survival modelling, and greater interpretability. These characteristics are particularly important in healthcare applications, where transparent reasoning, reliable risk stratification, and clinically interpretable predictions are essential for effective decision-making and personalised patient management (Rudin, 2019; Topol, 2019).

4.1.3 ROC Curve Comparison

Building on the previous comparison, model discrimination can be further evaluated using Receiver Operating Characteristic (ROC) analysis.

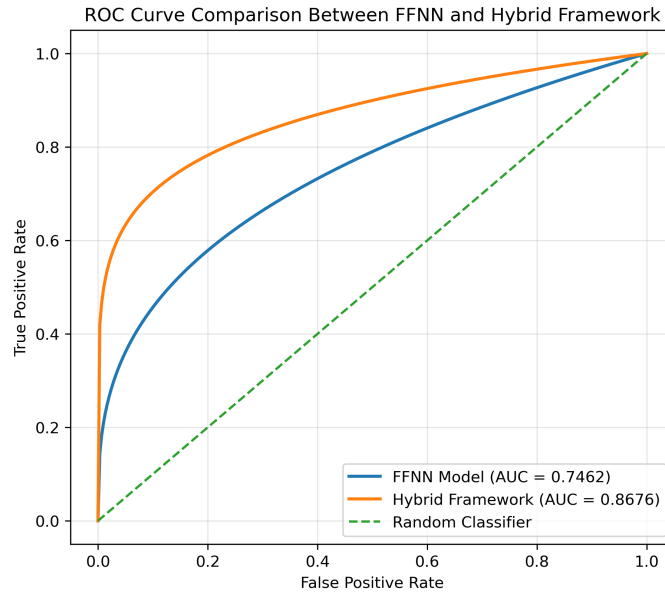


Figure 4.1: ROC curve comparison between the FFNN model and the proposed hybrid framework.

Figure 4.1 presents the ROC curves for the FFNN baseline model and the proposed hybrid framework. The ROC curve illustrates the relationship between the true positive rate (sensitivity) and the false positive rate (1-specificity) across different classification thresholds. Models with stronger discrimination performance produce curves closer to the upper-left corner of the graph, indicating better separation between positive and negative cases. The results show that the hybrid framework consistently outperforms the FFNN across most threshold values. This improvement is reflected in the higher AUC value of 0.8676 achieved by the hybrid framework, compared with 0.7462 for the FFNN model. The difference indicates that the proposed framework is more effective at separating patients with higher cardiovascular risk from lower-risk individuals. In particular, the hybrid framework maintains stronger discrimination at lower false positive rates, which is important in clinical settings where unnecessary interventions and false alarms should

be minimised. From a practical perspective, stronger ROC performance provides greater flexibility in clinical decision-making because risk thresholds may vary depending on the objective, such as screening, preventive monitoring, early intervention, or high-risk patient management. A model with more stable discrimination across thresholds therefore supports more reliable and adaptable clinical assessment. Overall, the ROC analysis further confirms that the proposed hybrid framework provides stronger and more consistent cardiovascular risk discrimination than the FFNN baseline model.

4.1.4 Interpretability Comparison

Table 4.2: Interpretability comparison between the FFNN model and the proposed hybrid framework

Aspect	FFNN	Hybrid Framework
Model Transparency	Low	High
Feature-Level Insight	Limited	Available
Risk Explanation	No	Yes (CPR-based)
Subgroup Analysis	No	Yes
Clinical Usability	Limited	Strong

Interpretability is an essential requirement in healthcare applications, where prediction outputs must be understandable, transparent, and clinically meaningful. Table 4.2 summarises the main interpretability differences between the FFNN baseline model and the proposed hybrid framework. The FFNN primarily operates as a black-box prediction model, where the relationship between input variables and prediction outcomes is not directly interpretable. Although the FFNN achieves reasonable predictive performance, it provides limited insight into how individual clinical variables contribute to cardiovascular risk. This lack of transparency may reduce clinician trust and limit practical clinical applicability. In contrast, the hybrid framework introduces interpretability at multiple stages of the modelling process. The clustering stage groups patients into clinically meaningful subgroups (see Section 3.5.1), allowing predictions to be interpreted within comparable patient profiles. The subgroup-specific Cox proportional hazards models then estimate the influence of clinical variables within each subgroup (see Section 3.4.2), providing additional insight into subgroup-level risk behaviour. Furthermore, the Cumulative Prevalence Ratio (CPR) summarises long-term cardiovascular risk in an interpretable form (see Section 3.4.3), making predictions easier to understand and communicate in clinical settings. Together, these components provide a structured and context-aware representation of cardiovascular risk. Overall, the hybrid framework achieves a stronger balance between predictive capability and interpretability, making it more suitable for practical clinical decision-making. The findings further suggest that the improvements introduced by the framework are not only numerical but also structural, as the model explicitly represents patient heterogeneity and temporal cardiovascular risk progression rather than relying on a single global prediction function.

4.1.5 Importance of Modelling Patient Heterogeneity

Following the discussion of interpretability, it is important to consider the role of patient heterogeneity in cardiovascular risk prediction. The results presented throughout Chapter 3 show that modelling heterogeneity is essential for meaningful cardiovascular risk assessment. Instead of assuming that all patients belong to a single homogeneous population, the proposed framework identifies clinically distinct subgroups through WTA-based clustering (see Section 3.5.1). These subgroups represent different patient profiles and are analysed separately during the survival modelling stage. The survival analysis results demonstrate that the identified subgroups follow different cardiovascular risk trajectories over time (see Section 3.5.2). This suggests that the relationship between clinical variables and cardiovascular outcomes is not uniform across the population. Rather, cardiovascular risk depends on complex interactions between patient characteristics that may vary substantially between subgroups. By explicitly modelling these differences, the framework provides a more clinically relevant and context-aware representation of cardiovascular risk. Although the clustering results show moderate subgroup separation (see Section 3.5.1), this behaviour is expected in real-world clinical datasets, where patient characteristics often vary continuously rather than forming perfectly distinct groups. The observed overlap therefore reflects realistic clinical complexity and suggests that the framework captures natural population variation instead of imposing artificially strict cluster boundaries. From a clinical perspective, subgroup-based modelling allows cardiovascular risk to be interpreted within comparable patient contexts, supporting more personalised and targeted assessment strategies. This subgroup-aware structure also provides the foundation for representing how cardiovascular risk changes over time.

4.1.6 Temporal Risk Representation and CPR

An important strength of the proposed framework is its ability to represent cardiovascular risk as a time-dependent process. Traditional classification models generally provide a single prediction at a fixed time point and therefore do not explicitly capture how cardiovascular risk evolves over time. In contrast, the proposed framework is developed within a survival-analysis setting, where cardiovascular risk is treated as a dynamic quantity that changes throughout the lifespan. The Cumulative Prevalence Ratio (CPR) further extends this interpretation by summarising cumulative cardiovascular risk across predefined time intervals (see Section 3.4.3). The CPR results are consistent with the subgroup survival patterns observed earlier (see Section 3.5.2), indicating that the framework successfully captures differences in long-term cardiovascular risk progression between patient subgroups. This provides a richer representation of patient outcomes because the framework models how cardiovascular risk accumulates over time rather than only predicting whether an event occurs. In addition, CPR enables direct comparison of long-term cardiovascular risk between patient subgroups, supporting more clinically meaningful interpretation and decision-making. Overall, the integration of subgroup-specific survival modelling with CPR-based cumulative risk estimation improves both interpretability and practical

usefulness by providing a clearer representation of cardiovascular risk progression within heterogeneous patient populations.

4.1.7 Clinical and Practical Implications

The findings of this study have important implications for clinical practice. In healthcare environments, decision-support systems must provide outputs that are not only accurate but also interpretable and clinically meaningful. The proposed hybrid framework addresses this requirement by linking predictions to subgroup structure and survival behaviour, allowing cardiovascular risk to be interpreted in relation to clinically similar patient profiles rather than isolated numerical outputs.

This improves transparency and supports more effective communication between clinicians and patients (Topol, 2019). In addition, the use of survival-based modelling aligns with established clinical methodologies commonly applied in epidemiology and medical risk assessment (Harrell, 2015). The subgroup-based structure of the framework also supports personalised cardiovascular risk assessment by recognising that different patient populations may follow different long-term risk trajectories. This is particularly important in cardiovascular disease, where risk progression is influenced by complex interactions between demographic, clinical, and lifestyle factors. From a practical perspective, the framework reduces reliance on post-hoc explanation techniques that are often required for black-box machine learning systems. Instead, interpretability is embedded directly within the modelling structure through clustering, subgroup-specific survival analysis, and CPR-based cumulative risk estimation. This improves transparency, reliability, and practical usability in clinical settings. Furthermore, the identification of higher-risk patient subgroups may support prioritisation of clinical interventions, screening strategies, and preventive management. Such subgroup-level reasoning may be particularly valuable in large-scale healthcare systems where efficient allocation of medical resources is important (Holzinger et al., 2017).

The ability to distinguish between different levels of long-term cardiovascular risk may also support earlier intervention and more targeted patient monitoring. Another important implication concerns the representation of cardiovascular risk over time. Traditional classification models generally provide a single static prediction outcome, whereas the proposed framework models cardiovascular risk as a dynamic process using survival analysis and cumulative risk estimation. This provides clinicians with a more informative understanding of long-term disease progression rather than only immediate event probability. The CPR-based representation further improves interpretability by summarising cumulative cardiovascular risk in a clinically understandable form. Overall, the proposed framework provides a balanced approach that combines predictive capability, subgroup-aware modelling, temporal risk representation, and clinical interpretability. These characteristics support its potential applicability in real-world healthcare environments and suggest that subgroup-based survival modelling may provide a valuable alternative to conventional global prediction

approaches.

4.1.8 Study-Level Evidence Positioning Using Published Cardiovascular Prediction Studies

To position the proposed hybrid framework within recent cardiovascular prediction research, a study-level evidence comparison was conducted using published studies that reported AUC or C-statistic values with confidence intervals where available. The purpose of this section is not to rank models directly or to claim that the proposed framework is universally superior. Instead, the aim is to show where the proposed framework is located within the current clinical prediction literature and to assess whether its discrimination performance is competitive with recent cardiovascular prediction models. Since published studies differ in population, endpoint definition, feature availability, follow-up period, and validation design, the comparison is interpreted as evidence positioning rather than a formal pooled analysis (Collins et al., 2024; Moons et al., 2015).

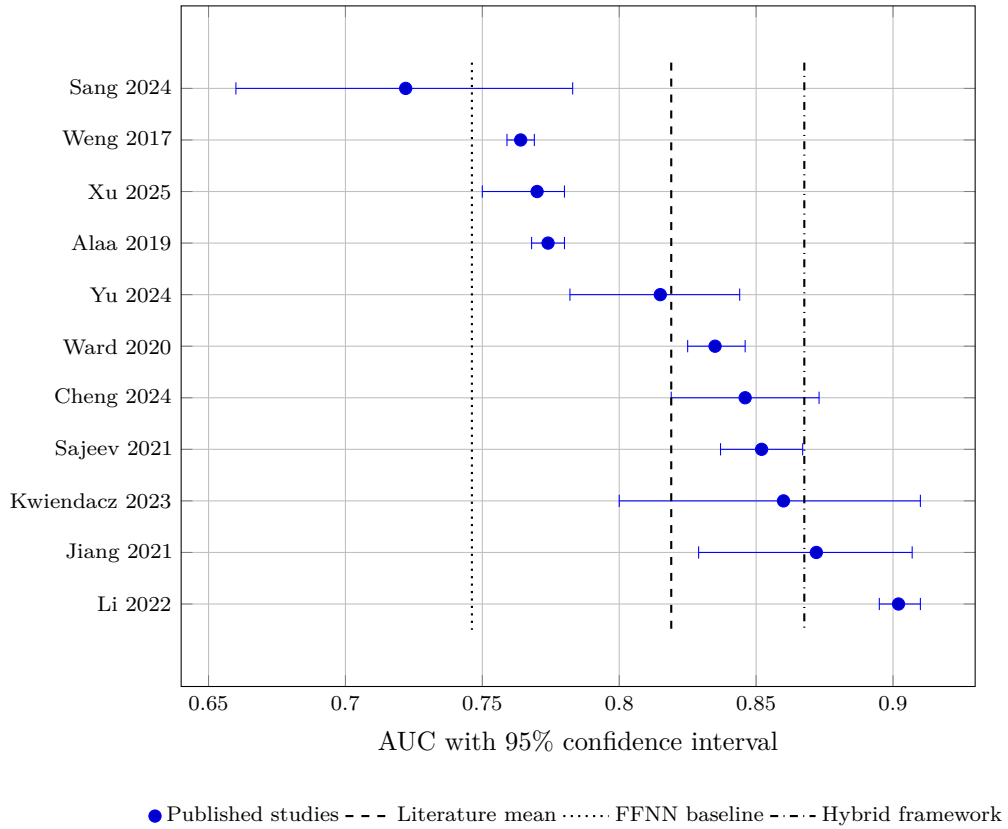


Figure 4.2: Forest-style evidence plot of cardiovascular prediction AUC values.

The comparison uses AUC-ROC because it is a standard discrimination measure in clinical prediction research. AUC measures how well a model distinguishes higher-risk patients from lower-risk patients across different classification thresholds (Fawcett, 2006). This makes it more suitable than accuracy for comparison across studies, because accuracy depends on one fixed threshold. Confidence intervals were included where available to show

the uncertainty around reported performance estimates. The selected studies cover a range of cardiovascular prediction settings, including ASCVD prediction, coronary artery disease prediction, cardiovascular mortality prediction, and diabetes-related cardiovascular risk prediction. They also include different modelling approaches, such as logistic regression, neural networks, ensemble models, and longitudinal deep learning methods. Studies with both moderate and high AUC values were included so that the proposed framework could be compared against the broader range of recent clinical prediction performance, including strong contemporary models. Figure 4.2 presents the AUC values from the selected studies together with the FFNN baseline and the proposed hybrid framework. Across the selected studies, AUC values range from approximately 0.722 to 0.902, with an approximate unweighted mean of 0.819. The proposed hybrid framework achieved an AUC of 0.8676 and a Harrell C-index of 0.8073, while the FFNN baseline achieved an AUC of 0.7462. The hybrid framework is therefore positioned above the literature mean and close to the upper-performing range of the selected studies. The AUC improvement of 0.1214 over the FFNN baseline suggests substantially stronger discrimination within the present dataset. Figure 4.2 provides an AUC-based overview of model positioning. However, AUC alone is not sufficient to judge clinical usefulness, because performance values depend on study design, data richness, validation strategy, and endpoint definition. Therefore, Table 4.3 provides additional context for the selected studies, including population type, endpoint, model structure, and key methodological considerations.

The studies summarised in Table 4.3 show that reported cardiovascular prediction performance varies across study settings, patient populations, endpoint definitions, predictor sets, and validation strategies. Within this context, the proposed hybrid framework achieved an AUC of 0.8676 together with a Harrell C-index of 0.8073, placing the framework within the upper-performing range of the selected contemporary cardiovascular prediction studies. The purpose of this evidence-positioning analysis is not to establish a direct ranking of models, but rather to examine how the proposed framework compares with recently published clinical prediction approaches. In addition to discrimination performance, the framework provides methodological advantages through WTA-based subgroup discovery, subgroup-specific survival modelling, and CPR-based cumulative cardiovascular risk estimation. Unlike conventional global prediction models, the proposed framework allows cardiovascular risk to be interpreted in relation to patient heterogeneity and temporal disease progression. Overall, the findings suggest that the framework performs competitively relative to recent cardiovascular prediction models while also maintaining interpretability and clinically meaningful subgroup-aware risk representation. Although further external validation and transportability assessment remain necessary before routine clinical deployment, the results support the value of combining predictive performance with structured and interpretable clinical modelling.

Table 4.3: Contextual comparison of published cardiovascular prediction studies used for evidence positioning.

Study	Population Endpoint	/ Model	AUC (95% CI)	Important Context and Limitations
Sang et al. (2024)	Korean hospital 3-year CVD prediction	T2DM cohorts; incident	Random Forest 0.722 (0.660–0.783)	External validation cohort contained only 32 CVD events. Internal AUC was higher (0.830), indicating reduced transportability. Diabetes-specific and hospital-based population limits broader applicability.
Weng et al. (2017)	UK primary care cohort; 10-year CVD prediction	Neural Network	0.764 (0.759–0.769)	Large real-world dataset with strong clinical relevance. However, validation used a same-source random split rather than independent external validation.
Xu et al. (2025)	NHANES population; prevalent CVD classification	T2DM XGBoost	0.770 (0.750–0.780)	Outcome was self-reported existing CVD rather than future incident events. Held-out validation AUC decreased to approximately 0.72.
Alaa et al. (2019)	UK Biobank participants; 5-year CVD prediction	AutoML	0.774 (0.768–0.780)	Model used 473 variables including questionnaire and non-traditional predictors. Strong performance partly reflects highly information-rich biobank data.
Yu et al. (2024)	Longitudinal ASCVD cohort; 10-year ASCVD prediction	Dynamic-DeepHi	0.815 (0.782–0.844)	Model required repeated longitudinal measurements over an 8-year observation period, increasing temporal feature richness.
Ward et al. (2020)	Northern California EHR cohort; 5-year ASCVD prediction	Gradient Boosting	0.835 (0.825–0.846)	Regional EHR-based cohort using structured healthcare records. Additional EHR variables produced only modest improvement beyond conventional predictors.
Cheng et al. (2024)	Taiwan Biobank cohort; CAD classification	Gradient Boosting	0.846 (0.819–0.873)	Outcome was self-reported CAD rather than prospective incident cardiovascular events. Logistic regression achieved similar performance (AUC 0.838).
Sajeev et al. (2021)	Australian cohorts; 15-year CVD mortality prediction	Machine Learning / LR	0.852 (0.837–0.867)	Predicted fatal CVD only rather than broader cardiovascular outcomes. Simpler linear models performed similarly to more complex methods.
Kwiendacz et al. (2023)	Single-centre diabetes cohort; cardiovascular phenotyping	Logistic Regression	0.860 (0.800–0.910)	Small cohort of 238 patients. Included treatment-related predictors that may reflect confounding by indication.
Jiang et al. (2021)	Kazakh population cohort; future CVD prediction	L1-regularised LR	0.872 (0.829–0.907)	Used biomarker-rich predictors including inflammatory cytokines. Validation relied on internal random splitting, and calibration showed risk overestimation.
Li et al. (2022)	Regional longitudinal EMR cohort; ASCVD prediction	RF-LTC	0.902 (0.895–0.910)	Highest reported AUC, but model used dense longitudinal EMR-derived laboratory and clinical features from a single healthcare system.

4.1.9 Summary of Discussion

This discussion examined the proposed hybrid cardiovascular risk framework from multiple perspectives, including predictive performance, subgroup modelling, interpretability, temporal risk representation, and evidence-based positioning within the existing literature. The findings show that the proposed framework improves discrimination performance compared with the FFNN baseline while also providing a more structured and clinically meaningful representation of cardiovascular risk. Unlike conventional global prediction models, the framework explicitly models patient heterogeneity through subgroup identification and cluster-specific survival analysis, allowing different patient profiles to exhibit distinct long-term cardiovascular risk trajectories. The integration of survival modelling and CPR-based cumulative risk estimation further enables cardiovascular risk to be represented as a dynamic and time-dependent process rather than as a single static prediction outcome. The evidence-based comparison with contemporary cardiovascular prediction studies also indicates that the proposed framework performs competitively relative to recent clinical prediction models while maintaining interpretability and subgroup-aware reasoning. Overall, the discussion demonstrates that the proposed framework achieves a balanced combination of predictive capability, interpretability, and clinical relevance. Nevertheless, additional external validation, calibration assessment, and evaluation across diverse populations remain necessary before routine clinical deployment can be considered.

4.2 Conclusion

This thesis addressed a central challenge in cardiovascular disease (CVD) risk prediction: achieving strong predictive performance while maintaining interpretability and clinical relevance. Traditional statistical models provide transparency and established clinical foundations but are often limited in their ability to capture complex and nonlinear relationships within heterogeneous patient populations. In contrast, modern machine learning approaches can achieve high predictive accuracy but frequently operate as black-box systems, reducing transparency and limiting practical clinical trust. To address this gap, this thesis proposed a hybrid cardiovascular risk prediction framework that integrates Winner-Takes-All (WTA) clustering, cluster-specific Cox proportional hazards modelling, and cumulative risk estimation through the Cumulative Prevalence Ratio (CPR). The proposed framework explicitly models patient heterogeneity by identifying data-driven subgroups and estimating risk separately within each subgroup. Unlike conventional global prediction models, this approach allows cardiovascular risk behaviour to vary across clinically similar patient groups. The WTA clustering mechanism adaptively identifies representative patient profiles without relying on predefined subgroup definitions, while the cluster-specific survival models capture subgroup-dependent relationships between clinical variables and cardiovascular outcomes. In addition, the CPR component provides an interpretable representation of cumulative long-term cardiovascular risk, allowing survival-based predictions to be translated into clinically meaningful risk categories.

The empirical results demonstrate that the proposed framework achieves strong predictive performance and reliable subgroup discrimination. The framework obtained a Harrell C-index of 0.8073 and an AUC of 0.8676, indicating strong discriminative capability and stable prediction performance. Comparative evaluation further showed that the proposed framework outperformed the FFNN baseline in discrimination performance while also providing substantially greater interpretability and subgroup-level reasoning. The survival analysis demonstrated that different patient subgroups follow distinct long-term risk trajectories, supporting the importance of modelling heterogeneity in cardiovascular prediction. Furthermore, calibration analysis and risk stratification evaluation confirmed that the framework produces clinically meaningful and reliable risk estimates. Beyond predictive performance, an important contribution of this thesis is the integration of interpretability directly within the modelling process. Rather than relying on post-hoc explanation techniques, the proposed framework provides transparent subgroup-level reasoning through clustering, survival modelling, and cumulative risk estimation. This enables clinicians to interpret predictions in relation to identifiable patient profiles and subgroup-specific survival behaviour. Such structured interpretability is particularly important in healthcare applications, where trust, transparency, and clinical understanding are essential for adoption and decision-making. The findings of this thesis therefore demonstrate that effective cardiovascular risk prediction requires more than high classification accuracy alone. Clinically meaningful prediction systems must also capture patient heterogeneity, model the temporal progression of disease risk, and provide interpretable representations of prediction behaviour. By integrating competitive learning, subgroup-based survival modelling, and cumulative risk estimation within a unified framework, this thesis presents a balanced and practically relevant approach to cardiovascular risk prediction. Overall, the proposed framework contributes to the advancement of interpretable artificial intelligence in healthcare and provides a strong methodological foundation for future research in heterogeneity-aware and time-dependent clinical risk modelling.

4.2.1 Future Directions

The findings of this thesis highlight several important directions for future research in cardiovascular disease (CVD) risk prediction, particularly in the development of hybrid, interpretable, and clinically deployable modelling frameworks. Although the proposed approach demonstrates strong predictive capability and interpretability, additional research is required to improve flexibility, robustness, and generalisability across different healthcare settings and patient populations. One important direction involves the integration of advanced representation learning techniques within interpretable hybrid architectures. While Feedforward Neural Networks (FFNNs) are capable of modelling nonlinear relationships, their limited interpretability restricts practical clinical adoption. Future work could therefore investigate hybrid architectures in which neural networks are used primarily for feature representation learning while preserving interpretable downstream

components such as subgroup clustering and survival-based modelling. Such approaches may improve predictive capability while maintaining transparency and clinically meaningful reasoning, consistent with recent developments in interpretable machine learning (Molnar, 2022; Rudin, 2019).

Another important extension concerns the survival modelling component of the framework. The current implementation relies on the Cox proportional hazards model, which assumes proportional hazard relationships over time. Although this assumption supports interpretability and stable estimation, it may not fully capture complex time-varying risk behaviour. Future research could therefore explore more flexible survival-analysis approaches, including time-dependent covariates, stratified Cox models, neural survival models, and deep hazard-based learning methods. Such extensions may allow more accurate modelling of dynamic cardiovascular risk progression and changing patient conditions over time (Bzdok et al., 2018; Katzman et al., 2018). The integration of multi-modal and high-dimensional healthcare data also represents a promising research direction. Modern clinical environments increasingly generate diverse forms of patient information, including genomic data, medical imaging, wearable sensor measurements, and longitudinal electronic health record (EHR) data. Extending the proposed framework to incorporate multi-modal information may improve predictive capability and support more personalised cardiovascular risk assessment. In particular, subgroup discovery across combined feature spaces may allow the identification of more complex and clinically informative patient phenotypes (Topol, 2019; Willeminck et al., 2020).

From a methodological perspective, the subgroup identification stage can also be further refined. Although the WTA competitive learning mechanism demonstrates adaptive behaviour and effective subgroup representation, alternative clustering strategies such as density-based methods, hierarchical clustering, and probabilistic mixture models may provide additional flexibility and robustness. Future work should also investigate subgroup stability analysis and uncertainty quantification to improve the reliability and reproducibility of subgroup identification (Bishop, 2006; Hastie et al., 2009; Rousseeuw, 1987). External validation and real-world clinical deployment remain critical future objectives. The current study is based on a single dataset, and evaluation across independent populations and healthcare systems is necessary to assess transportability and generalisability. Multi-centre validation studies and prospective clinical evaluation would provide stronger evidence regarding practical utility and reliability. In addition, integration within electronic health record systems may support real-time cardiovascular risk prediction and early intervention strategies. Close collaboration with clinicians will remain essential to ensure that future systems maintain both interpretability and practical clinical usability (Topol, 2019).

Finally, ethical considerations must remain central to future developments in clinical artificial intelligence. Issues related to fairness, bias, transparency, and patient privacy are particularly important in healthcare prediction systems. Biases within training data may lead to unequal predictive performance across demographic or clinical subgroups, potentially reinforcing existing healthcare disparities. Future research should therefore

include fairness-aware evaluation, bias mitigation strategies, and continuous monitoring of deployed systems (Obermeyer et al., 2019). In addition, privacy-preserving approaches such as federated learning may provide opportunities for large-scale collaborative model development while maintaining regulatory compliance and patient confidentiality (Willemink et al., 2020).

In summary, future research should focus on extending the proposed framework through improved survival modelling, integration of advanced representation learning, incorporation of multi-modal healthcare data, robust external validation, and ethical clinical deployment. Maintaining a balance between predictive performance, interpretability, and responsible artificial intelligence will remain essential for the development of reliable and clinically applicable cardiovascular risk prediction systems.

References

- Akoglu, H. (2018). User’s guide to correlation coefficients. *Turkish Journal of Emergency Medicine*, 18(3), 91–93. <https://doi.org/10.1016/j.tjem.2018.08.001>
- Alaa, A. M., Bolton, T., Di Angelantonio, E., Rudd, J. H. F., & van der Schaar, M. (2019). Cardiovascular disease risk prediction using automated machine learning: A prospective study of 423,604 UK biobank participants. *PLOS ONE*, 14(5), e0213653. <https://doi.org/10.1371/journal.pone.0213653>
- American Diabetes Association. (2022). Standards of medical care in diabetes—2022. *Diabetes Care*, 45(Supplement₁), S1–S264. <https://doi.org/10.2337/dc22-SINT>
- Benjamin, E. J., Virani, S. S., Callaway, C. W., Chamberlain, A. M., Chang, A. R., Cheng, S., Chiuve, S. E., Cushman, M., Delling, F. N., Deo, R., et al. (2017). Heart disease and stroke statistics—2017 update. *Circulation*, 135(10), e146–e603. <https://doi.org/10.1161/CIR.0000000000000485>
- Bishop, C. M. (2006). *Pattern recognition and machine learning*. Springer. <https://doi.org/10.1007/978-0-387-45528-0>
- Bzdok, D., Altman, N., & Krzywinski, M. (2018). Statistics versus machine learning. *Nature Methods*, 15(4), 233–234. <https://doi.org/10.1038/nmeth.4642>
- Chen, T., & Guestrin, C. (2016). Xgboost: A scalable tree boosting system. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 785–794. <https://doi.org/10.1145/2939672.2939785>
- Cheng, C.-H., Lee, B.-J., Nfor, O. N., et al. (2024). Using machine learning-based algorithms to construct cardiovascular risk prediction models for taiwanese adults based on traditional and novel risk factors. *BMC Medical Informatics and Decision Making*, 24, 199. <https://doi.org/10.1186/s12911-024-02603-2>
- Ching, T., Himmelstein, D. S., Beaulieu-Jones, B. K., Kalinin, A. A., Do, B. T., Way, G. P., Ferrero, E., Agapow, P.-M., Zietz, M., Hoffman, M. M., et al. (2018). Opportunities and obstacles for deep learning in biology and medicine. *Journal of the Royal Society Interface*, 15(141). <https://doi.org/10.1098/rsif.2017.0387>
- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Routledge. <https://doi.org/10.4324/9780203771587>
- Collins, G. S., Dhiman, P., Andaur Navarro, C. L., Ma, J., Hooft, L., Reitsma, J. B., Logullo, P., Beam, A. L., Peng, L., Van Calster, B., & Moons, K. G. M. (2024). TRIPOD+AI statement: Updated guidance for reporting clinical prediction models that use regression or machine learning methods. *BMJ*, 385, e078378. <https://doi.org/10.1136/bmj-2023-078378>
- Cox, D. R. (1972). Regression models and life-tables. *Journal of the Royal Statistical Society: Series B*, 34(2), 187–220. <https://doi.org/10.1111/j.2517-6161.1972.tb00899.x>

- DeFilippis, A. P., Young, R., Carrubba, C. J., McEvoy, J. W., Budoff, M. J., Blumenthal, R. S., Kronmal, R. A., Criqui, M. H., Nambi, V., et al. (2015). An analysis of calibration and discrimination among multiple cardiovascular risk scores in a modern multiethnic cohort. *Annals of Internal Medicine*, 162(4), 266–275. <https://doi.org/10.7326/M14-1281>
- Detrano, R., Janosi, A., Steinbrunn, W., et al. (1989). International application of a new probability algorithm for the diagnosis of coronary artery disease. *American Journal of Cardiology*, 64(5), 304–310. [https://doi.org/10.1016/0002-9149\(89\)90524-9](https://doi.org/10.1016/0002-9149(89)90524-9)
- Doshi-Velez, F., & Kim, B. (2017). Towards a rigorous science of interpretable machine learning. *arXiv preprint arXiv:1702.08608*. <https://doi.org/10.48550/arXiv.1702.08608>
- Esteva, A., Robicquet, A., Ramsundar, B., et al. (2019). A guide to deep learning in healthcare. *Nature Medicine*, 25(1), 24–29. <https://doi.org/10.1038/s41591-018-0316-z>
- Fawcett, T. (2006). An introduction to roc analysis. *Pattern Recognition Letters*, 27(8), 861–874. <https://doi.org/10.1016/j.patrec.2005.10.010>
- Ference, B. A., Ginsberg, H. N., Graham, I., Ray, K. K., Packard, C. J., Bruckert, E., Hegele, R. A., Krauss, R. M., Raal, F. J., Schunkert, H., Watts, G. F., Borén, J., Fazio, S., Horton, J. D., Masana, L., Nicholls, S. J., Nordestgaard, B. G., van de Sluis, B., Taskinen, M. R., ... Catapano, A. L. (2017). Low-density lipoproteins cause atherosclerotic cardiovascular disease. *European Heart Journal*, 38(32), 2459–2472. <https://doi.org/10.1093/eurheartj/ehx144>
- Goodfellow, I., Bengio, Y., & Courville, A. (2016). *Deep learning*. MIT Press. <https://www.deeplearningbook.org/>
- Harrell, F. E. (2015). *Regression modeling strategies: With applications to linear models, logistic regression, and survival analysis* (2nd ed.). Springer. <https://doi.org/10.1007/978-3-319-19425-7>
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). *The elements of statistical learning: Data mining, inference, and prediction* (2nd ed.). Springer. <https://doi.org/10.1007/978-0-387-84858-7>
- Haykin, S. (1999). *Neural networks: A comprehensive foundation*. Prentice Hall.
- Hippisley-Cox, J., & Coupland, C. (2017). Development and validation of qrisk3 risk prediction algorithms to estimate future risk of cardiovascular disease: Prospective cohort study. *BMJ*, 357, j2099. <https://doi.org/10.1136/bmj.j2099>
- Hochreiter, S., & Schmidhuber, J. (1997). Long short-term memory. *Neural Computation*, 9(8), 1735–1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
- Holzinger, A., Biemann, C., Pattichis, C. S., & Kell, D. B. (2017). What do we need to build explainable AI systems for the medical domain? <https://arxiv.org/abs/1712.09923>
- Hosmer, D. W., Lemeshow, S., & May, S. (2008). *Applied survival analysis: Regression modeling of time-to-event data* (2nd ed.). Wiley. <https://doi.org/10.1002/9780470258019>

- Hruby, A., & Hu, F. B. (2015). The epidemiology of obesity: A big picture. *Circulation Research*, 117(3), 277–291. <https://doi.org/10.1161/CIRCRESAHA.115.306883>
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2021). *An introduction to statistical learning* (2nd ed.). Springer. <https://doi.org/10.1007/978-1-0716-1418-1>
- Jiang, Y., Zhang, X., Ma, R., et al. (2021). Cardiovascular disease prediction by machine learning algorithms based on cytokines in kazakhs of china. *Clinical Epidemiology*, 13, 417–428. <https://doi.org/10.2147/CLEP.S313343>
- Kaplan, E. L., & Meier, P. (1958). Nonparametric estimation from incomplete observations. *Journal of the American Statistical Association*, 53(282), 457–481. <https://doi.org/10.1080/01621459.1958.10501452>
- Katzman, J. L., Shaham, U., Bates, J., Cloninger, A., Jiang, T., & Kluger, Y. (2018). Deepsurv: Personalized treatment recommender system using a cox proportional hazards deep neural network. *BMC Medical Research Methodology*, 18(1), 24. <https://doi.org/10.1186/s12874-018-0482-1>
- Kingma, D. P., & Ba, J. (2015). Adam: A method for stochastic optimization. *International Conference on Learning Representations (ICLR)*. <https://doi.org/10.48550/arXiv.1412.6980>
- Komorowski, M., Celi, L. A., Badawi, O., Gordon, A. C., & Faisal, A. R. (2018). The artificial intelligence clinician learns optimal treatment strategies for sepsis in intensive care. *Nature Medicine*, 24, 1716–1720. <https://doi.org/10.1038/s41591-018-0213-5>
- Krittanawong, C., Johnson, K. W., & Rosenson, R. S. (2020). Machine learning prediction in cardiovascular diseases: A meta-analysis. *Scientific Reports*, 10, 16057. <https://doi.org/10.1038/s41598-020-72685-1>
- Kröse, B., & van der Smagt, P. (1996). An introduction to neural networks [Available online]. <https://www.infor.uva.es/~teodoro/neuro-intro.pdf>
- Kwiendacz, H., Wijata, A. M., Nalepa, J., Piaśnik, J., Kulpa, J., Herba, M., et al. (2023). Machine learning profiles of cardiovascular risk in patients with diabetes mellitus: The silesia diabetes-heart project. *Cardiovascular Diabetology*, 22, 218. <https://doi.org/10.1186/s12933-023-01938-w>
- LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436–444. <https://doi.org/10.1038/nature14539>
- Li, Q., Campan, A., Ren, A., & Eid, W. E. (2022). Automating and improving cardiovascular disease prediction using machine learning and EMR data features from a regional healthcare system. *International Journal of Medical Informatics*, 163, 104786. <https://doi.org/10.1016/j.ijmedinf.2022.104786>
- Liyanage, H., & Lipnicka, M. (2026). Computational framework for dynamic cardiovascular risk assessment with cluster-specific cox models and cumulative risk analysis. *Archives of Control Sciences*, 36(1), 99–131. <https://doi.org/10.24425/acs.2026.158423>
- McCulloch, W. S., & Pitts, W. (1943). A logical calculus of the ideas immanent in nervous activity. *Bulletin of Mathematical Biophysics*, 5, 115–133. <https://doi.org/10.1007/BF02478259>

- Minsky, M., & Papert, S. (1969). *Perceptrons: An introduction to computational geometry*. MIT Press. <https://mitpress.mit.edu/9780262630221/perceptrons/>
- Molnar, C. (2022). *Interpretable machine learning* (2nd ed.). Lulu.com. <https://christophm.github.io/interpretable-ml-book/>
- Moons, K. G. M., Altman, D. G., Reitsma, J. B., Ioannidis, J. P. A., Macaskill, P., Steyerberg, E. W., Vickers, A. J., Ransohoff, D. F., & Collins, G. S. (2015). Transparent reporting of a multivariable prediction model for individual prognosis or diagnosis (TRIPOD): Explanation and elaboration. *Annals of Internal Medicine*, 162(1), W1–W73. <https://doi.org/10.7326/M14-0698>
- Nair, V., & Hinton, G. E. (2010). Rectified linear units improve restricted boltzmann machines. *Proceedings of the 27th International Conference on Machine Learning (ICML)*, 807–814. <https://www.cs.toronto.edu/~fritz/absps/reluICML.pdf>
- Ng, M., Fleming, T., Robinson, M., Thomson, B., Graetz, N., Margono, C., Mullany, E. C., Biryukov, S., Abbafati, C., Abera, S. F., et al. (2014). Global, regional, and national prevalence of overweight and obesity in children and adults during 1980–2013: A systematic analysis for the global burden of disease study 2013. *The Lancet*, 384(9945), 766–781. [https://doi.org/10.1016/S0140-6736\(14\)60460-8](https://doi.org/10.1016/S0140-6736(14)60460-8)
- Novikoff, A. B. J. (1962). On convergence proofs for perceptrons. *Proceedings of the Symposium on the Mathematical Theory of Automata*, 615–622.
- Obermeyer, Z., Powers, B., Vogeli, C., & Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464), 447–453. <https://doi.org/10.1126/science.aax2342>
- Prechelt, L. (1998). Early stopping — but when? In *Neural networks: Tricks of the trade* (pp. 55–69). Springer. https://doi.org/10.1007/3-540-49430-8_3
- Rosenblatt, F. (1958). The perceptron: A probabilistic model for information storage and organization in the brain. *Psychological Review*, 65(6), 386–408. <https://doi.org/10.1037/h0042519>
- Rousseeuw, P. J. (1987). Silhouettes: A graphical aid to the interpretation and validation of cluster analysis. *Journal of Computational and Applied Mathematics*, 20, 53–65. [https://doi.org/10.1016/0377-0427\(87\)90125-7](https://doi.org/10.1016/0377-0427(87)90125-7)
- Rudin, C. (2019). Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. *Nature Machine Intelligence*, 1, 206–215. <https://doi.org/10.1038/s42256-019-0048-x>
- Rumelhart, D. E., Hinton, G. E., & Williams, R. J. (1986). Learning representations by back-propagating errors. *Nature*, 323, 533–536. <https://doi.org/10.1038/323533a0>
- Sajeev, S., Champion, S., Beleigoli, A., et al. (2021). Predicting australian adults at high risk of cardiovascular disease mortality using standard risk factors and machine learning. *International Journal of Environmental Research and Public Health*, 18(6), 3187. <https://doi.org/10.3390/ijerph18063187>
- Sang, H., Lee, H., Lee, M., Park, J., Kim, S., Woo, H. G., Rahmati, M., et al. (2024). Prediction model for cardiovascular disease in patients with diabetes using machine

- learning derived and validated in two independent korean cohorts. *Scientific Reports*, 14, 14966. <https://doi.org/10.1038/s41598-024-63798-y>
- Shwartz-Ziv, R., & Armon, A. (2022). Tabular data: Deep learning is not all you need. *Information Fusion*, 81, 84–90. <https://doi.org/10.1016/j.inffus.2021.11.011>
- Silverman, B. W. (1986). *Density estimation for statistics and data analysis*. Chapman; Hall. <https://doi.org/10.1007/978-1-4899-3324-9>
- Sokolova, M., & Lapalme, G. (2009). A systematic analysis of performance measures for classification tasks. *Information Processing & Management*, 45(4), 427–437. <https://doi.org/10.1016/j.ipm.2009.03.002>
- Therneau, T. M., & Grambsch, P. M. (2000). *Modeling survival data: Extending the cox model*. Springer. <https://doi.org/10.1007/978-1-4757-3294-8>
- Topol, E. J. (2019). High-performance medicine: The convergence of human and artificial intelligence. *Nature Medicine*, 25(1), 44–56. <https://doi.org/10.1038/s41591-018-0300-7>
- Tukey, J. W. (1977). *Exploratory data analysis*. Addison-Wesley.
- Warburton, D. E. R., Nicol, C. W., & Bredin, S. S. D. (2006). Health benefits of physical activity: The evidence. *CMAJ*, 174(6), 801–809. <https://doi.org/10.1503/cmaj.051351>
- Ward, A., Sarraju, A., Chung, S., Li, J., Harrington, R., Heidenreich, P., Palaniappan, L., & Scheinker, D. (2020). Machine learning and atherosclerotic cardiovascular disease risk prediction in a multi-ethnic population. *npj Digital Medicine*, 3, 125. <https://doi.org/10.1038/s41746-020-00331-1>
- Weng, S. F., Reps, J., Kai, J., Garibaldi, J. M., & Qureshi, N. (2017). Can machine-learning improve cardiovascular risk prediction using routine clinical data? *PLOS ONE*, 12(4), e0174944. <https://doi.org/10.1371/journal.pone.0174944>
- Willemink, M. J., Koszek, W. A., Hardell, C., et al. (2020). Preparing medical imaging data for machine learning. *Radiology*, 295(1), 4–15. <https://doi.org/10.1148/radiol.2020192224>
- Wilson, P. W. F., D’Agostino, R. B., Levy, D., Belanger, A. M., Silbershatz, H., & Kannel, W. B. (1998). Prediction of coronary heart disease using risk factor categories. *Circulation*, 97(18), 1837–1847. <https://doi.org/10.1161/01.CIR.97.18.1837>
- World Health Organization. (2021). Cardiovascular diseases (cvds). [https://www.who.int/news-room/fact-sheets/detail/cardiovascular-diseases-\(cvds\)](https://www.who.int/news-room/fact-sheets/detail/cardiovascular-diseases-(cvds))
- Xu, C., Shi, F., Ding, W., Fang, C., & Fang, C. (2025). Development and validation of a machine learning model for cardiovascular disease risk prediction in type 2 diabetes patients. *Scientific Reports*, 15, 32818. <https://doi.org/10.1038/s41598-025-18443-7>
- Yu, J., Yang, X., Deng, Y., Krefman, A. E., Pool, L. R., Zhao, L., et al. (2024). Incorporating longitudinal history of risk factors into atherosclerotic cardiovascular disease risk prediction using deep learning. *Scientific Reports*, 14, 2554. <https://doi.org/10.1038/s41598-024-51685-5>

Appendix

Appendix 1: Feedforward Neural Network Implementation

This appendix presents the Python implementation used for the Feedforward Neural Network (FFNN) model described in Chapter 2. The implementation includes data loading, preprocessing, architecture comparison, model training, performance evaluation, and visualisation of results.

Listing 1: General implementation of the FFNN-based cardiovascular risk prediction model

```
1 import os
2 import random
3 import numpy as np
4 import pandas as pd
5 import matplotlib.pyplot as plt
6 import tensorflow as tf
7
8 from sklearn.model_selection import train_test_split
9 from sklearn.preprocessing import MinMaxScaler
10 from sklearn.metrics import (
11     accuracy_score, precision_score, recall_score,
12     f1_score, roc_auc_score, confusion_matrix,
13     classification_report, ConfusionMatrixDisplay,
14     RocCurveDisplay
15 )
16
17 from tensorflow.keras.models import Sequential
18 from tensorflow.keras.layers import Dense, Input
19 from tensorflow.keras.optimizers import Adam
20 from tensorflow.keras.callbacks import EarlyStopping
21
22
23 # -----
24 # 0. Reproducibility
25 # -----
26 os.environ["PYTHONHASHSEED"] = "42"
27 random.seed(42)
28 np.random.seed(42)
29 tf.random.set_seed(42)
30
31
32 # -----
33 # 1. Load dataset
34 # -----
35 df = pd.read_csv("dataset.csv")
36
37
38 # -----
39 # 2. Features and target variable
40 # -----
41 features = [
42     "Age", "Gender", "BMI", "Diastolic_Blood_Pressure",
```

```

43     "Cholesterol_Level", "Smoke", "Alcohol",
44     "Active", "Fasting_Blood_Sugar"
45 ]
46
47 target = "CVD"
48
49 X = df[features]
50 y = df[target]
51
52
53 # -----
54 # 3. Train-test split
55 # -----
56 X_train, X_test, y_train, y_test = train_test_split(
57     X, y,
58     test_size=0.20,
59     random_state=42,
60     stratify=y
61 )
62
63
64 # -----
65 # 4. Feature scaling
66 # -----
67 scaler = MinMaxScaler()
68 X_train = scaler.fit_transform(X_train)
69 X_test = scaler.transform(X_test)
70
71
72 # -----
73 # 5. Candidate FFNN architectures
74 # -----
75 architectures = [
76     [10, 20, 10, 1],
77     [15, 30, 15, 1],
78     [20, 50, 20, 1],
79     [25, 60, 25, 1],
80     [30, 70, 30, 1],
81     [35, 80, 35, 1],
82     [40, 90, 40, 1],
83     [45, 100, 45, 1],
84     [50, 110, 50, 1],
85     [60, 120, 60, 1]
86 ]
87
88
89 # -----
90 # 6. FFNN model builder
91 # -----
92 def build_model(arch):
93     model = Sequential()
94     model.add(Input(shape=(9,)))
95
96     for units in arch[:-1]:
97         model.add(Dense(units, activation="relu"))
98
99     model.add(Dense(1, activation="sigmoid"))
100
101     model.compile(

```

```

102         optimizer=Adam(learning_rate=0.001),
103         loss="binary_crossentropy",
104         metrics=["accuracy"]
105     )
106
107     return model
108
109
110 # -----
111 # 7. Train and evaluate architectures
112 # -----
113 results = []
114 histories = {}
115 models = {}
116
117 early_stop = EarlyStopping(
118     monitor="val_loss",
119     patience=10,
120     restore_best_weights=True
121 )
122
123 for arch in architectures:
124     model = build_model(arch)
125
126     history = model.fit(
127         X_train,
128         y_train,
129         epochs=200,
130         batch_size=64,
131         validation_split=0.20,
132         callbacks=[early_stop],
133         verbose=0
134     )
135
136     y_prob = model.predict(X_test, verbose=0).ravel()
137     y_pred = (y_prob >= 0.5).astype(int)
138
139     results.append({
140         "Architecture": str(arch),
141         "Accuracy": accuracy_score(y_test, y_pred),
142         "Precision": precision_score(y_test, y_pred, zero_division=0),
143         "Recall": recall_score(y_test, y_pred, zero_division=0),
144         "F1-score": f1_score(y_test, y_pred, zero_division=0),
145         "AUC": roc_auc_score(y_test, y_prob),
146         "Epochs_used": len(history.history["loss"])
147     })
148
149     histories[str(arch)] = history
150     models[str(arch)] = model
151
152
153 # -----
154 # 8. Select best model
155 # -----
156 results_df = pd.DataFrame(results).sort_values(by="Accuracy", ascending=False)
157
158 best_arch_str = results_df.iloc[0]["Architecture"]
159 best_model = models[best_arch_str]
160 best_history = histories[best_arch_str]

```

```
161
162
163 # -----
164 # 9. Final evaluation
165 # -----
166 y_prob = best_model.predict(X_test, verbose=0).ravel()
167 y_pred = (y_prob >= 0.5).astype(int)
168
169 cm = confusion_matrix(y_test, y_pred)
170
171
172 # -----
173 # 10. Visualisations
174 # -----
175 ConfusionMatrixDisplay(confusion_matrix=cm).plot()
176 plt.savefig("confusion_matrix.png", dpi=300)
177
178 RocCurveDisplay.from_predictions(y_test, y_prob)
179 plt.savefig("roc_curve.png", dpi=300)
180
181 plt.plot(best_history.history["loss"], label="Train")
182 plt.plot(best_history.history["val_loss"], label="Validation")
183 plt.legend()
184 plt.savefig("loss_curve.png", dpi=300)
```

Hybrid Framework Implementation

Appendix 2: This appendix presents the general Python implementation used for the hybrid computational framework described in Chapter 3. The code includes data preprocessing, Winner-Takes-All (WTA) clustering, cluster-specific Cox proportional hazards modelling, CPR-based risk estimation, and model evaluation.

Listing 2: General implementation of the hybrid WTA–Cox–CPR framework

```

1 import os
2 import random
3 import warnings
4 import numpy as np
5 import pandas as pd
6 import matplotlib.pyplot as plt
7
8 from sklearn.preprocessing import MinMaxScaler
9 from sklearn.decomposition import PCA
10 from sklearn.manifold import TSNE
11 from sklearn.metrics import (
12     silhouette_score,
13     davies_bouldin_score,
14     calinski_harabasz_score,
15     roc_auc_score
16 )
17 from sklearn.model_selection import StratifiedKFold
18 from sklearn.linear_model import LogisticRegression
19 from statsmodels.duration.hazard_regression import PHReg
20
21 warnings.filterwarnings("ignore")
22
23 # -----
24 # 1. Configuration
25 # -----
26 DATA_PATH = "CVD_SL Dataset.csv"
27 OUTPUT_DIR = "hybrid_results"
28
29 os.makedirs(OUTPUT_DIR, exist_ok=True)
30
31 RANDOM_STATE = 42
32 np.random.seed(RANDOM_STATE)
33 random.seed(RANDOM_STATE)
34
35 N_CLUSTERS = 15
36 WTA_EPOCHS = 40
37 WTA_LEARNING_RATE = 0.05
38 WTA_MIN_LEARNING_RATE = 0.005
39
40 CPR_TA = 20.0
41 CPR_TB = 99.0
42
43 TIME_COL = "Age"
44 EVENT_COL = "CVD"
45
46 FEATURES = [
47     "Age",
48     "Gender",
49     "BMI",

```

```

50     "Diastolic_Blood_Pressure",
51     "Cholesterol_Level",
52     "Smoke",
53     "Alcohol",
54     "Active",
55     "Fasting_Blood_Sugar"
56 ]
57
58 CLUSTERING_FEATURES = FEATURES
59
60 COX_COVARIATES = [
61     "Gender",
62     "BMI",
63     "Diastolic_Blood_Pressure",
64     "Cholesterol_Level",
65     "Smoke",
66     "Alcohol",
67     "Active",
68     "Fasting_Blood_Sugar"
69 ]
70
71
72 # -----
73 # 2. Helper functions
74 # -----
75 def ensure_numeric(df, columns):
76     df = df.copy()
77     for col in columns:
78         df[col] = pd.to_numeric(df[col], errors="coerce")
79     return df
80
81
82 def recode_categorical_variables(df):
83     df = df.copy()
84
85     if "Gender" in df.columns:
86         unique_gender = set(df["Gender"].dropna().unique())
87         if unique_gender.issubset({1, 2}):
88             df["Gender"] = df["Gender"].map({1: 1, 2: 0})
89
90     return df
91
92
93 def minmax_scale(df, columns):
94     df_scaled = df.copy()
95     scaler = MinMaxScaler()
96     df_scaled[columns] = scaler.fit_transform(df_scaled[columns])
97     return df_scaled, scaler
98
99
100 class WTAClusterer:
101     def __init__(
102         self,
103         n_units=15,
104         epochs=40,
105         learning_rate=0.05,
106         min_learning_rate=0.005,
107         random_state=42
108     ):

```

```

109     self.n_units = n_units
110     self.epochs = epochs
111     self.learning_rate = learning_rate
112     self.min_learning_rate = min_learning_rate
113     self.random_state = random_state
114     self.weights_ = None
115     self.history_ = []
116
117     def fit(self, X):
118         rng = np.random.default_rng(self.random_state)
119         X = np.asarray(X, dtype=float)
120
121         n_samples, n_features = X.shape
122         init_idx = rng.choice(n_samples, size=self.n_units, replace=False)
123         W = X[init_idx].copy()
124
125         for epoch in range(self.epochs):
126             eta = self.learning_rate * (1 - epoch / max(self.epochs - 1, 1))
127             eta = max(eta, self.min_learning_rate)
128
129             order = rng.permutation(n_samples)
130
131             for i in order:
132                 x = X[i]
133                 activations = W @ x
134                 winner = np.argmax(activations)
135                 W[winner] = W[winner] + eta * (x - W[winner])
136
137             labels = np.argmax(X @ W.T, axis=1)
138
139             self.history_.append({
140                 "Epoch": epoch + 1,
141                 "Learning_Rate": eta,
142                 "Active_Clusters": len(np.unique(labels))
143             })
144
145             self.weights_ = W
146             return self
147
148         def predict(self, X):
149             X = np.asarray(X, dtype=float)
150             return np.argmax(X @ self.weights_.T, axis=1)
151
152         def fit_predict(self, X):
153             self.fit(X)
154             return self.predict(X)
155
156
157     def breslow_baseline_hazard(times, events, linear_predictor):
158         times = np.asarray(times, dtype=float)
159         events = np.asarray(events, dtype=int)
160         linear_predictor = np.asarray(linear_predictor, dtype=float)
161
162         event_times = np.sort(np.unique(times[events == 1]))
163         exp_lp = np.exp(np.clip(linear_predictor, -50, 50))
164
165         increments = []
166
167         for t in event_times:

```

```

168     d_i = np.sum((times == t) & (events == 1))
169     risk_sum = np.sum(exp_lp[times >= t])
170
171     if risk_sum > 0:
172         increments.append(d_i / risk_sum)
173     else:
174         increments.append(np.nan)
175
176     increments = np.asarray(increments)
177     cumulative_hazard = np.nancumsum(increments)
178
179     return event_times, increments, cumulative_hazard
180
181
182 def compute_cpr(event_times, baseline_increments, eta_cluster, ta, tb):
183     mask = (event_times >= ta) & (event_times <= tb)
184
185     if np.sum(mask) == 0:
186         return np.nan
187
188     cumulative_baseline_hazard = np.nansum(baseline_increments[mask])
189     cpr = np.exp(np.clip(eta_cluster, -50, 50)) * cumulative_baseline_hazard
190
191     return float(cpr)
192
193
194 def harrell_c_index(times, events, risk_scores):
195     concordant = 0.0
196     tied = 0.0
197     comparable = 0.0
198
199     n = len(times)
200
201     for i in range(n):
202         for j in range(i + 1, n):
203             if times[i] == times[j]:
204                 continue
205
206             if times[i] < times[j] and events[i] == 1:
207                 comparable += 1
208                 if risk_scores[i] > risk_scores[j]:
209                     concordant += 1
210                 elif risk_scores[i] == risk_scores[j]:
211                     tied += 1
212
213             elif times[j] < times[i] and events[j] == 1:
214                 comparable += 1
215                 if risk_scores[j] > risk_scores[i]:
216                     concordant += 1
217                 elif risk_scores[j] == risk_scores[i]:
218                     tied += 1
219
220     if comparable == 0:
221         return np.nan
222
223     return (concordant + 0.5 * tied) / comparable
224
225
226 # -----

```

```

227 # 3. Load and preprocess dataset
228 # -----
229 df = pd.read_csv(DATA_PATH)
230
231 df = ensure_numeric(df, FEATURES + [EVENT_COL])
232 df = recode_categorical_variables(df)
233
234 df = df.dropna(subset=FEATURES + [EVENT_COL])
235 df = df[df[EVENT_COL].isin([0, 1])]
236 df = df[df[TIME_COL] > 0]
237
238 df_raw = df.copy()
239
240 df_scaled, scaler = minmax_scale(df, CLUSTERING_FEATURES)
241
242 X_cluster = df_scaled[CLUSTERING_FEATURES].values
243 y = df_raw[EVENT_COL].astype(int)
244
245
246 # -----
247 # 4. WTA clustering
248 # -----
249 wta = WTAClusterer(
250     n_units=N_CLUSTERS,
251     epochs=WTA_EPOCHS,
252     learning_rate=WTA_LEARNING_RATE,
253     min_learning_rate=WTA_MIN_LEARNING_RATE,
254     random_state=RANDOM_STATE
255 )
256
257 cluster_labels = wta.fit_predict(X_cluster)
258
259 df_scaled["Cluster"] = cluster_labels + 1
260 df_raw["Cluster"] = cluster_labels + 1
261
262 wta_history = pd.DataFrame(wta.history_)
263 wta_history.to_csv(
264     os.path.join(OUTPUT_DIR, "wta_training_history.csv"),
265     index=False
266 )
267
268 cluster_sizes = df_raw["Cluster"].value_counts().sort_index().reset_index()
269 cluster_sizes.columns = ["Cluster", "Number_of_Patients"]
270 cluster_sizes["Percentage"] = (
271     cluster_sizes["Number_of_Patients"] /
272     cluster_sizes["Number_of_Patients"].sum() * 100
273 )
274
275 cluster_sizes.to_csv(
276     os.path.join(OUTPUT_DIR, "cluster_size_distribution.csv"),
277     index=False
278 )
279
280
281 # -----
282 # 5. Clustering validation
283 # -----
284 silhouette = silhouette_score(X_cluster, cluster_labels)
285 davies_bouldin = davies_bouldin_score(X_cluster, cluster_labels)

```

```

286 calinski_harabasz = calinski_harabasz_score(X_cluster, cluster_labels)
287
288 cluster_validation = pd.DataFrame({
289     "Metric": [
290         "Silhouette",
291         "Davies_Bouldin",
292         "Calinski_Harabasz"
293     ],
294     "Value": [
295         silhouette,
296         davies_bouldin,
297         calinski_harabasz
298     ]
299 })
300
301 cluster_validation.to_csv(
302     os.path.join(OUTPUT_DIR, "cluster_validation_metrics.csv"),
303     index=False
304 )
305
306
307 # -----
308 # 6. PCA and t-SNE visualisation
309 # -----
310 pca = PCA(n_components=2, random_state=RANDOM_STATE)
311 pca_data = pca.fit_transform(X_cluster)
312
313 plt.figure(figsize=(8, 6))
314 plt.scatter(
315     pca_data[:, 0],
316     pca_data[:, 1],
317     c=df_raw["Cluster"],
318     s=5,
319     alpha=0.7
320 )
321 plt.title("PCA Visualisation of WTA Clusters")
322 plt.xlabel("PCA Component 1")
323 plt.ylabel("PCA Component 2")
324 plt.tight_layout()
325 plt.savefig(
326     os.path.join(OUTPUT_DIR, "patient_cluster_visualization_pca.png"),
327     dpi=300
328 )
329 plt.close()
330
331 sample_size = min(4000, len(df_scaled))
332 sample_df = df_scaled.sample(sample_size, random_state=RANDOM_STATE)
333
334 tsne = TSNE(
335     n_components=2,
336     random_state=RANDOM_STATE,
337     init="pca",
338     learning_rate="auto"
339 )
340
341 tsne_data = tsne.fit_transform(sample_df[CLUSTERING_FEATURES])
342
343 plt.figure(figsize=(8, 6))
344 plt.scatter(

```

```

345     tsne_data[:, 0],
346     tsne_data[:, 1],
347     c=sample_df["Cluster"],
348     s=5,
349     alpha=0.7
350 )
351 plt.title("t-SNE Visualisation of WTA Clusters")
352 plt.xlabel("t-SNE Component 1")
353 plt.ylabel("t-SNE Component 2")
354 plt.tight_layout()
355 plt.savefig(
356     os.path.join(OUTPUT_DIR, "patient_cluster_visualization_tsne.png"),
357     dpi=300
358 )
359 plt.close()
360
361
362 # -----
363 # 7. Cluster-specific Cox modelling and CPR computation
364 # -----
365 df_survival = df_raw.copy()
366 df_survival["time"] = df_survival[TIME_COL].astype(float)
367 df_survival["event"] = df_survival[EVENT_COL].astype(int)
368
369 cox_status_rows = []
370 cox_coefficient_rows = []
371 hazard_ratio_rows = []
372 cpr_rows = []
373
374 patient_risk_scores = np.full(len(df_survival), np.nan)
375
376 for cluster_id in sorted(df_survival["Cluster"].unique()):
377     cluster_data = df_survival[df_survival["Cluster"] == cluster_id].copy()
378
379     X_cox = cluster_data[COX_COVARIATES].copy()
380     times = cluster_data["time"].values
381     events = cluster_data["event"].values
382
383     if len(cluster_data) < 50 or np.sum(events) < 15:
384         cox_status_rows.append([
385             cluster_id,
386             len(cluster_data),
387             int(np.sum(events)),
388             "Skipped"
389         ])
390         continue
391
392     try:
393         model = PHReg(
394             endog=times,
395             exog=X_cox,
396             status=events
397         )
398
399         result = model.fit(dispen=0)
400
401         beta = np.asarray(result.params, dtype=float)
402         linear_predictor = X_cox.values @ beta
403

```

```

404     patient_risk_scores[cluster_data.index] = linear_predictor
405
406     cox_status_rows.append([
407         cluster_id,
408         len(cluster_data),
409         int(np.sum(events)),
410         "Fitted"
411     ])
412
413     for var, coef in zip(X_cox.columns, beta):
414         cox_coefficient_rows.append([
415             cluster_id,
416             var,
417             coef
418         ])
419
420         hazard_ratio_rows.append([
421             cluster_id,
422             var,
423             np.exp(coef)
424         ])
425
426     event_times, baseline_increments, cumulative_hazard = (
427         breslow_baseline_hazard(
428             times,
429             events,
430             linear_predictor
431         )
432     )
433
434     cluster_centroid = X_cox.mean().values
435     eta_cluster = float(cluster_centroid @ beta)
436
437     cpr_value = compute_cpr(
438         event_times,
439         baseline_increments,
440         eta_cluster,
441         CPR_TA,
442         CPR_TB
443     )
444
445     cpr_rows.append([
446         cluster_id,
447         cpr_value
448     ])
449
450     except Exception:
451         cox_status_rows.append([
452             cluster_id,
453             len(cluster_data),
454             int(np.sum(events)),
455             "Skipped"
456         ])
457
458
459 cox_status = pd.DataFrame(
460     cox_status_rows,
461     columns=[
462         "Cluster",

```

```

463         "Cluster_Size",
464         "Number_of_Events",
465         "Cox_Status"
466     ]
467 )
468
469 cox_status.to_csv(
470     os.path.join(OUTPUT_DIR, "clusterwise_cox_fit_status.csv"),
471     index=False
472 )
473
474 cox_coefficients = pd.DataFrame(
475     cox_coefficient_rows,
476     columns=[
477         "Cluster",
478         "Variable",
479         "Coefficient"
480     ]
481 )
482
483 cox_coefficients.to_csv(
484     os.path.join(OUTPUT_DIR, "cluster_specific_cox_coefficients.csv"),
485     index=False
486 )
487
488 hazard_ratios = pd.DataFrame(
489     hazard_ratio_rows,
490     columns=[
491         "Cluster",
492         "Variable",
493         "Hazard_Ratio"
494     ]
495 )
496
497 hazard_ratios.to_csv(
498     os.path.join(OUTPUT_DIR, "hazard_ratios.csv"),
499     index=False
500 )
501
502 cluster_cpr = pd.DataFrame(
503     cpr_rows,
504     columns=[
505         "Cluster",
506         "CPR_Value"
507     ]
508 )
509
510 cluster_cpr.to_csv(
511     os.path.join(OUTPUT_DIR, "cluster_cpr_values.csv"),
512     index=False
513 )
514
515
516 # -----
517 # 8. CPR-based risk classification
518 # -----
519 q20, q50, q80 = cluster_cpr["CPR_Value"].quantile([0.2, 0.5, 0.8])
520
521 def assign_risk_level(cpr):

```

```

522     if cpr <= q20:
523         return "Low"
524     elif cpr <= q50:
525         return "Moderate"
526     elif cpr <= q80:
527         return "High"
528     else:
529         return "Critical"
530
531 cluster_cpr["Risk_Level"] = cluster_cpr["CPR_Value"].apply(assign_risk_level)
532
533 cluster_cpr.to_csv(
534     os.path.join(OUTPUT_DIR, "cluster_cpr_risk_levels.csv"),
535     index=False
536 )
537
538 cpr_thresholds = pd.DataFrame({
539     "Risk_Category": ["Low", "Moderate", "High", "Critical"],
540     "Threshold_Rule": [
541         f"CPR <= {q20:.6f}",
542         f"{q20:.6f} < CPR <= {q50:.6f}",
543         f"{q50:.6f} < CPR <= {q80:.6f}",
544         f"CPR > {q80:.6f}"
545     ]
546 })
547
548 cpr_thresholds.to_csv(
549     os.path.join(OUTPUT_DIR, "cpr_thresholds.csv"),
550     index=False
551 )
552
553
554 # -----
555 # 9. Harrell C-index
556 # -----
557 valid_mask = np.isfinite(patient_risk_scores)
558
559 c_index = harrell_c_index(
560     df_survival.loc[valid_mask, "time"].values,
561     df_survival.loc[valid_mask, "event"].values,
562     patient_risk_scores[valid_mask]
563 )
564
565 c_index_df = pd.DataFrame({
566     "Metric": ["Harrell_C_Index"],
567     "Value": [c_index]
568 })
569
570 c_index_df.to_csv(
571     os.path.join(OUTPUT_DIR, "c_index_performance.csv"),
572     index=False
573 )
574
575
576 # -----
577 # 10. Auxiliary classification evaluation
578 # -----
579 cv = StratifiedKfold(
580     n_splits=5,

```

```

581     shuffle=True,
582     random_state=RANDOM_STATE
583 )
584
585 cv_results = []
586
587 for fold, (train_idx, test_idx) in enumerate(
588     cv.split(df_scaled[CLUSTERING_FEATURES], y),
589     start=1
590 ):
591     X_train = df_scaled.iloc[train_idx][CLUSTERING_FEATURES]
592     X_test = df_scaled.iloc[test_idx][CLUSTERING_FEATURES]
593     y_train = y.iloc[train_idx]
594     y_test = y.iloc[test_idx]
595
596     clf = LogisticRegression(
597         max_iter=2000,
598         random_state=RANDOM_STATE
599     )
600
601     clf.fit(X_train, y_train)
602     y_prob = clf.predict_proba(X_test)[:, 1]
603
604     auc = roc_auc_score(y_test, y_prob)
605
606     cv_results.append([
607         fold,
608         auc
609     ])
610
611 cv_results_df = pd.DataFrame(
612     cv_results,
613     columns=["Fold", "AUC"]
614 )
615
616 cv_results_df.to_csv(
617     os.path.join(OUTPUT_DIR, "auxiliary_auc_results.csv"),
618     index=False
619 )
620
621
622 # -----
623 # 11. Final summary
624 # -----
625 summary_metrics = pd.DataFrame({
626     "Metric": [
627         "Silhouette",
628         "Davies_Bouldin",
629         "Calinski_Harabasz",
630         "Harrell_C_Index",
631         "Auxiliary_Mean_AUC",
632         "Auxiliary_Std_AUC"
633     ],
634     "Value": [
635         silhouette,
636         davies_bouldin,
637         calinski_harabasz,
638         c_index,
639         cv_results_df["AUC"].mean(),

```

```
640         cv_results_df["AUC"].std()
641     ]
642 })
643
644 summary_metrics.to_csv(
645     os.path.join(OUTPUT_DIR, "summary_metrics.csv"),
646     index=False
647 )
648
649 print("Hybrid WTA-Cox-CPR analysis completed successfully.")
650 print("Results saved in:", OUTPUT_DIR)
```

Appendix 3: Prediction Output and Risk Category Analysis

This appendix presents the additional Python implementation used to generate prediction-level outputs from the final hybrid WTA–Cox–CPR framework. The code summarises CPR-based risk categories, cluster-level prediction patterns, observed CVD rates by predicted risk level, and representative patient examples. The script is written in a general format and should be executed after the main hybrid framework code.

Listing 3: Prediction output and CPR-based risk category analysis

```

1 import os
2 import pandas as pd
3 import numpy as np
4
5 # -----
6 # 1. Configuration
7 # -----
8 DATA_PATH = "dataset.csv"
9 BASE_DIR = os.path.dirname(DATA_PATH)
10
11 OUTPUT_DIR = "hybrid_results"
12 PREDICTION_DIR = os.path.join(OUTPUT_DIR, "prediction_results")
13
14 os.makedirs(PREDICTION_DIR, exist_ok=True)
15
16 analysis_path = os.path.join(
17     OUTPUT_DIR,
18     "analysis_dataset_with_wta_clusters_and_risk.csv"
19 )
20
21 if not os.path.exists(analysis_path):
22     raise FileNotFoundError(
23         "Run the main WTA--Cox--CPR framework first."
24     )
25
26
27 # -----
28 # 2. Load final analysis dataset
29 # -----
30 pred_df = pd.read_csv(analysis_path)
31
32 prediction_df = pred_df.dropna(
33     subset=["Cluster", "CPR_Value", "Risk_Level"]
34 ).copy()
35
36
37 # -----
38 # 3. Predicted risk category distribution
39 # -----
40 risk_distribution = (
41     prediction_df["Risk_Level"]
42     .value_counts()
43     .reindex(["Low", "Moderate", "High", "Critical"])
44     .fillna(0)
45     .reset_index()
46 )
47
48 risk_distribution.columns = [

```

```

49     "Predicted_Risk_Category",
50     "Number_of_Patients"
51 ]
52
53 risk_distribution["Percentage"] = (
54     risk_distribution["Number_of_Patients"]
55     / risk_distribution["Number_of_Patients"].sum()
56     * 100
57 ).round(2)
58
59 risk_distribution.to_csv(
60     os.path.join(PREDICTION_DIR, "predicted_risk_category_distribution.csv"),
61     index=False
62 )
63
64
65 # -----
66 # 4. Cluster-level prediction summary
67 # -----
68 cluster_prediction_summary = (
69     prediction_df
70     .groupby(["Cluster", "Risk_Level", "CPR_Value"])
71     .size()
72     .reset_index(name="Number_of_Patients")
73     .sort_values("Cluster")
74 )
75
76 cluster_prediction_summary["Percentage"] = (
77     cluster_prediction_summary["Number_of_Patients"]
78     / cluster_prediction_summary["Number_of_Patients"].sum()
79     * 100
80 ).round(2)
81
82 cluster_prediction_summary.to_csv(
83     os.path.join(PREDICTION_DIR, "cluster_level_prediction_summary.csv"),
84     index=False
85 )
86
87
88 # -----
89 # 5. Observed CVD rate by predicted risk level
90 # -----
91 if "CVD" in prediction_df.columns:
92     observed_summary = (
93         prediction_df
94         .groupby("Risk_Level")["CVD"]
95         .agg(
96             Number_of_Patients="count",
97             Observed_CVD_Cases="sum",
98             Observed_CVD_Rate="mean"
99         )
100         .reindex(["Low", "Moderate", "High", "Critical"])
101         .reset_index()
102     )
103
104     observed_summary["Observed_CVD_Rate"] = (
105         observed_summary["Observed_CVD_Rate"] * 100
106     ).round(2)
107

```

```

108     observed_summary.to_csv(
109         os.path.join(PREDICTION_DIR, "observed_cvd_rate_by_predicted_risk.csv"),
110         index=False
111     )
112
113
114 # -----
115 # 6. Representative patient examples
116 # -----
117 example_cols = [
118     "Age_Raw",
119     "Gender",
120     "BMI",
121     "Diastolic_Blood_Pressure",
122     "Cholesterol_Level",
123     "Smoke",
124     "Alcohol",
125     "Active",
126     "Fasting_Blood_Sugar",
127     "Cluster",
128     "CPR_Value",
129     "Risk_Level"
130 ]
131
132 available_example_cols = [
133     col for col in example_cols if col in prediction_df.columns
134 ]
135
136 representative_examples = (
137     prediction_df
138     .groupby("Risk_Level", group_keys=False)
139     .apply(lambda x: x.sample(min(3, len(x)), random_state=42))
140     [available_example_cols]
141     .reset_index(drop=True)
142 )
143
144 representative_examples.to_csv(
145     os.path.join(PREDICTION_DIR, "representative_patient_prediction_examples.csv"),
146     index=False
147 )
148
149
150 # -----
151 # 7. Save LaTeX tables
152 # -----
153 risk_distribution.to_latex(
154     os.path.join(PREDICTION_DIR, "predicted_risk_category_distribution.tex"),
155     index=False,
156     caption="Distribution of predicted CPR-based risk categories.",
157     label="tab:prediction_distribution"
158 )
159
160 cluster_prediction_summary.to_latex(
161     os.path.join(PREDICTION_DIR, "cluster_level_prediction_summary.tex"),
162     index=False,
163     caption="Cluster-level prediction summary from the hybrid framework.",
164     label="tab:cluster_prediction_summary"
165 )
166

```

```
167 if "CVD" in prediction_df.columns:
168     observed_summary.to_latex(
169         os.path.join(PREDICTION_DIR, "observed_cvd_rate_by_predicted_risk.tex"),
170         index=False,
171         caption="Observed CVD rate by predicted CPR-based risk category.",
172         label="tab:observed_cvd_by_risk"
173     )
174
175 representative_examples.to_latex(
176     os.path.join(PREDICTION_DIR, "representative_patient_prediction_examples.tex"),
177     index=False,
178     caption="Representative patient-level prediction examples.",
179     label="tab:representative_prediction_examples"
180 )
181
182 print("Prediction result tables saved in:", PREDICTION_DIR)
```

Appendix 4: Chapter 3 Hybrid Framework Evaluation Code

This appendix presents the Python implementation used to generate the Chapter 3 evaluation outputs for the proposed hybrid cardiovascular risk framework. The code produces clustering evaluation, survival analysis, CPR-based risk stratification, prediction analysis, and visualisation outputs generated from the implemented WTA–Cox–CPR framework. The script should be executed after the main hybrid framework pipeline has produced the final clustering assignments, survival estimates, and CPR values.

Listing 4: Chapter 3 hybrid framework evaluation and prediction analysis

```

1 import os
2 import numpy as np
3 import pandas as pd
4 import matplotlib.pyplot as plt
5
6 from sklearn.metrics import (
7     silhouette_score,
8     davies_bouldin_score,
9     calinski_harabasz_score,
10    roc_auc_score
11 )
12
13 from sklearn.decomposition import PCA
14 from sklearn.manifold import TSNE
15
16 from lifelines import KaplanMeierFitter
17
18
19 # -----
20 # 1. Configuration
21 # -----
22 OUTPUT_DIR = "hybrid_results"
23
24 TABLE_DIR = os.path.join(OUTPUT_DIR, "tables")
25 FIGURE_DIR = os.path.join(OUTPUT_DIR, "figures")
26
27 os.makedirs(TABLE_DIR, exist_ok=True)
28 os.makedirs(FIGURE_DIR, exist_ok=True)
29
30 analysis_path = os.path.join(
31     OUTPUT_DIR,
32     "analysis_dataset_with_wta_clusters_and_risk.csv"
33 )
34
35 if not os.path.exists(analysis_path):
36     raise FileNotFoundError(
37         "Run the main hybrid framework before evaluation."
38     )
39
40 df = pd.read_csv(analysis_path)
41
42
43 # -----
44 # 2. Required columns
45 # -----
46 required_cols = [

```

```

47     "Cluster",
48     "Age",
49     "CVD",
50     "CPR",
51     "Cox_Risk_Score_LP"
52 ]
53
54 missing_cols = [
55     col for col in required_cols
56     if col not in df.columns
57 ]
58
59 if missing_cols:
60     raise ValueError(
61         "Missing required columns: "
62         + str(missing_cols)
63     )
64
65
66 # -----
67 # 3. Clustering quality metrics
68 # -----
69 feature_cols = [
70     col for col in df.columns
71     if col not in [
72         "Cluster",
73         "Age",
74         "CVD",
75         "CPR",
76         "Cox_Risk_Score_LP"
77     ]
78 ]
79
80 X = df[feature_cols].values
81 cluster_labels = df["Cluster"].values
82
83 sil_score = silhouette_score(X, cluster_labels)
84
85 db_score = davies_bouldin_score(
86     X,
87     cluster_labels
88 )
89
90 ch_score = calinski_harabasz_score(
91     X,
92     cluster_labels
93 )
94
95 cluster_validation = pd.DataFrame([
96     {
97         "Silhouette_Score": sil_score,
98         "Davies_Bouldin_Index": db_score,
99         "Calinski_Harabasz_Index": ch_score
100     }
101 ])
102
103 cluster_validation.to_csv(
104     os.path.join(TABLE_DIR, "cluster_validation.csv"),
105     index=False

```

```

106 )
107
108
109 # -----
110 # 4. PCA visualisation
111 # -----
112 pca = PCA(n_components=2)
113
114 X_pca = pca.fit_transform(X)
115
116 plt.figure(figsize=(7, 6))
117
118 scatter = plt.scatter(
119     X_pca[:, 0],
120     X_pca[:, 1],
121     c=cluster_labels,
122     s=10
123 )
124
125 plt.xlabel("Principal Component 1")
126 plt.ylabel("Principal Component 2")
127 plt.title("PCA-Based Visualisation of WTA Clusters")
128
129 plt.tight_layout()
130
131 plt.savefig(
132     os.path.join(FIGURE_DIR, "patient_cluster_visualization_pca.png"),
133     dpi=300,
134     bbox_inches="tight"
135 )
136
137 plt.close()
138
139
140 # -----
141 # 5. t-SNE visualisation
142 # -----
143 tsne = TSNE(
144     n_components=2,
145     perplexity=30,
146     random_state=42
147 )
148
149 X_tsne = tsne.fit_transform(X)
150
151 plt.figure(figsize=(7, 6))
152
153 scatter = plt.scatter(
154     X_tsne[:, 0],
155     X_tsne[:, 1],
156     c=cluster_labels,
157     s=10
158 )
159
160 plt.xlabel("t-SNE Dimension 1")
161 plt.ylabel("t-SNE Dimension 2")
162 plt.title("t-SNE Visualisation of WTA Clusters")
163
164 plt.tight_layout()

```

```

165
166 plt.savefig(
167     os.path.join(FIGURE_DIR, "patient_cluster_visualization_tsne.png"),
168     dpi=300,
169     bbox_inches="tight"
170 )
171
172 plt.close()
173
174
175 # -----
176 # 6. Kaplan--Meier survival analysis
177 # -----
178 kmf = KaplanMeierFitter()
179
180 plt.figure(figsize=(8, 6))
181
182 selected_clusters = [1, 4, 8, 13]
183
184 for cluster_id in selected_clusters:
185
186     cluster_df = df[
187         df["Cluster"] == cluster_id
188     ]
189
190     kmf.fit(
191         durations=cluster_df["Age"],
192         event_observed=cluster_df["CVD"],
193         label=f"Cluster {cluster_id}"
194     )
195
196     kmf.plot_survival_function()
197
198 plt.xlabel("Age")
199 plt.ylabel("Survival Probability")
200
201 plt.title(
202     "Kaplan--Meier Survival Curves"
203 )
204
205 plt.grid(alpha=0.3)
206
207 plt.tight_layout()
208
209 plt.savefig(
210     os.path.join(FIGURE_DIR, "kaplan_meier_survival_curves.png"),
211     dpi=300,
212     bbox_inches="tight"
213 )
214
215 plt.close()
216
217
218 # -----
219 # 7. CPR-based risk categorisation
220 # -----
221 risk_thresholds = {
222     "Low": 103.537827,
223     "Moderate": 130.928023,

```

```

224     "High": 146.094782
225 }
226
227 def assign_risk(cpr):
228
229     if cpr <= risk_thresholds["Low"]:
230         return "Low"
231
232     elif cpr <= risk_thresholds["Moderate"]:
233         return "Moderate"
234
235     elif cpr <= risk_thresholds["High"]:
236         return "High"
237
238     else:
239         return "Critical"
240
241
242 df["Risk_Category"] = df["CPR"].apply(
243     assign_risk
244 )
245
246 risk_distribution = (
247     df["Risk_Category"]
248     .value_counts()
249     .reset_index()
250 )
251
252 risk_distribution.columns = [
253     "Risk_Category",
254     "Patients"
255 ]
256
257 risk_distribution.to_csv(
258     os.path.join(TABLE_DIR, "risk_distribution.csv"),
259     index=False
260 )
261
262
263 # -----
264 # 8. Prediction analysis
265 # -----
266 prediction_summary = (
267     df.groupby("Cluster")
268     .agg(
269         Patients=("Cluster", "size"),
270         Mean_CPR=("CPR", "mean"),
271         Event_Rate=("CVD", "mean")
272     )
273     .reset_index()
274 )
275
276 prediction_summary.to_csv(
277     os.path.join(TABLE_DIR, "prediction_summary.csv"),
278     index=False
279 )
280
281
282 # -----

```

```

283 # 9. ROC analysis
284 # -----
285 valid_mask = np.isfinite(
286     df["Cox_Risk_Score_LP"]
287 )
288
289 y_true = df.loc[
290     valid_mask,
291     "CVD"
292 ].astype(int).values
293
294 risk_score = df.loc[
295     valid_mask,
296     "Cox_Risk_Score_LP"
297 ].values
298
299 auc_value = roc_auc_score(
300     y_true,
301     risk_score
302 )
303
304 auc_table = pd.DataFrame([
305     {
306         "AUC": auc_value
307     }
308 ])
309
310 auc_table.to_csv(
311     os.path.join(TABLE_DIR, "auc_results.csv"),
312     index=False
313 )
314
315
316 # -----
317 # 10. Save LaTeX tables
318 # -----
319 cluster_validation.round(4).to_latex(
320     os.path.join(TABLE_DIR, "cluster_validation.tex"),
321     index=False,
322     caption="Clustering validation metrics.",
323     label="tab:cluster_validation"
324 )
325
326 risk_distribution.to_latex(
327     os.path.join(TABLE_DIR, "risk_distribution.tex"),
328     index=False,
329     caption="Distribution of predicted risk categories.",
330     label="tab:risk_distribution"
331 )
332
333 prediction_summary.round(4).to_latex(
334     os.path.join(TABLE_DIR, "prediction_summary.tex"),
335     index=False,
336     caption="Cluster-level prediction summary.",
337     label="tab:prediction_summary"
338 )
339
340 auc_table.round(4).to_latex(
341     os.path.join(TABLE_DIR, "auc_results.tex"),

```

```
342     index=False,
343     caption="Auxiliary AUC evaluation of the hybrid framework.",
344     label="tab:auc_results"
345 )
346
347 print("Chapter 3 evaluation outputs generated successfully.")
348 print("Tables saved to:", TABLE_DIR)
349 print("Figures saved to:", FIGURE_DIR)
```