



**Politecnico
di Torino**

Politecnico di Torino

Master Degree in Electronic Engineering

A.a. 2025/2026

Graduation Session July 2026

Automated Detection of Brugada Syndrome Type 3 from 12-Lead ECGs Using a Three-Stream Deep Learning Architecture

Supervisor:

Prof. Pasero Eros Gian Alessandro

Candidate:

Muhammad Hammad Rauf

Co-supervisor:

Randazzo Vincenzo

Abstract

Brugada Syndrome (BrS) is an inherited cardiac channelopathy associated with sudden arrhythmic death in structurally normal hearts. Among its three recognised ECG pattern types, the Type-3 pattern is defined by right precordial ST elevation below 2 mm in leads V1–V2 and is the most morphologically subtle and clinically ambiguous manifestation. No published machine learning pipeline has previously addressed the binary discrimination of Type-3 BrS from non-Brugada (Non-BrS) ECGs as a dedicated classification task under a patient-disjoint evaluation protocol.

This thesis presents BrugadaECGNet, an end-to-end pipeline for Type-3 BrS classification from standard 12-lead ECGs. A dataset of 237 recordings was assembled from two heterogeneous acquisition sources: 63 Type-3 BrS recordings in GE Sapphire DCAR XML format from a clinical centre in Torino, Italy, and 174 Non-BrS recordings digitised from paper ECGs. A source-confound audit identified a 40-fold pre-QRS baseline variance disparity between the two source populations. Per-channel baseline normalisation reduced this ratio to 1.84, with z-score ablation confirming that residual classification is driven by morphological signal rather than acquisition fingerprint. All results are reported under 5-fold patient-disjoint cross-validation, in which all recordings from the same patient are confined to a single fold.

BrugadaECGNet combines three information streams: (A) a PTB-XL pretrained CNN-Transformer encoder with a custom hierarchical attention layer operating on eight precordial channels; (B) a multi-layer perceptron over a 44-dimensional clinically grounded morphological feature vector incorporating the de Luna α/β -angle descriptors, the Corrado index, the Crea inferior ST-depression criterion, the Babai-Bigi aVR sign, and the r-J interval in V1; and (C) a shallow convolutional encoder over eight inferior-lead channels capturing the temporal profile of inferior ST morphology. The three streams are fused in a shared classification head with a Type-3 F1-optimising threshold (T3-opt).

Across five patient-disjoint folds, the pipeline achieves a recording-level ROC-AUC (RecAUC) of 0.939 ± 0.032 , recording-level accuracy (RecAcc) of 0.872 ± 0.018 , Non-BrS F1 (RecNBF1) of 0.911 ± 0.016 , and Type-3 BrS F1 (RecT3F1) of 0.766 ± 0.034 at the T3-opt threshold. These results are competitive with

the published ML benchmark range (AUC 0.93–0.97) and exceed documented cardiologist performance (AUC 0.72–0.88). A systematic post-baseline intervention study which are comprising of test-time augmentation, multi-seed ensembling, and feature vector expansion, established that the RecT3F1 ceiling is a data-level constraint imposed by the limited number of unique Type-3 patients (≈ 15 –20 in the current dataset), not an architectural limitation. The source-confound audit methodology and the patient-disjoint evaluation framework are transferable to any multi-source ECG classification study.

Acknowledgements

I wish to express my sincere gratitude to the people who have supported me throughout the course of this thesis. First and foremost, I extend my thanks to my supervisor, Professor Pasero Eros Gian Alessandro and Co-supervisor Randazzo Vioncenzo, for providing me with the opportunity to pursue this research and for his continuous support, insightful feedback, and availability, all of which have been very valuable for the successful completion of this thesis.

I am especially grateful to my thesis mentor, Alessandro for his dedicated supervision, constant availability, patience, and technical guidance. This work would not have been possible without his constant support and encouragement.

Finally, on a personal note I would like to thank Allah, my wife, my parents, my entire family, and my friends for their collective love, support, encouragement, and prayers. Their belief in me has been a constant source of strength and motivation throughout the course of this work.

Table of Contents

List of Tables	IX
List of Figures	XIII
List of Acronyms	XVI
1 Introduction	1
1.1 Background and Clinical Motivation	1
1.2 Limitations of Current ECG Diagnosis	2
1.3 Machine Learning for Brugada Syndrome: The Gap	3
1.4 Contributions of This Work	3
1.5 Thesis Outline	4
2 Background and Literature Review	5
2.1 Brugada Syndrome — Clinical and Genetic Overview	5
2.2 ECG Morphology of Brugada Patterns	5
2.3 Prior Machine Learning Work for Brugada Classification	8
2.3.1 CNN-Based Approaches	9
2.3.2 Deep Learning with Transfer Learning	9
2.3.3 Benchmarks and the Patient-Disjoint Gap	9
2.3.4 Risk Stratification and Clinical Feature Studies	9
2.4 Attention Mechanisms and Transformer Architectures for ECG	11
2.5 Transfer Learning and Domain Adaptation in ECG	12
2.6 Multi-Source ECG Dataset Challenges and Confound Mitigation . .	13
3 Dataset and Data Preparation	15
3.1 Digitized Paper ECGs	15
3.2 Type-3 BrS Class: GE Sapphire DCAR XML	15
3.3 Clinical Labelling and Dataset Curation	16

3.4	Patient Grouping and Cross-Validation Design	17
3.5	Source-Confound Audit	18
4	Pipeline Development	21
4.1	Design Principles and Evaluation Criteria	21
4.2	Phase 1 — Preprocessing Infrastructure and Dataset Integration . .	22
4.3	Phase 2 — Architecture Development: From BiLSTM to Transformer	22
4.4	Phase 3 — PTB-XL Transfer Learning	23
4.5	Phase 4 — Hand-Engineered Feature Engineering (44-Dimensional Vector)	25
4.6	Phase 5 — Recording-Level Aggregation and Threshold Calibration	29
4.7	Phase 6 — Three-Stream Architecture: Stream C Addition	30
4.8	Summary — Architecture Evolution and Ablation Results	33
5	Final Model Architecture and Training	35
5.1	Input Representation	35
5.2	BrugadaECGNet Architecture	39
5.3	Training Protocol	41
5.4	Recording-Level Inference and Threshold Calibration	43
5.5	Feature Imputation and Scaling	44
6	Experimental Results and Performance Evaluation	45
6.1	Cross-Validation Performance Overview	45
6.2	ROC Curve Analysis	47
6.3	Confusion Matrix Analysis	50
6.4	Threshold Analysis	52
6.5	Attention Weight Analysis	52
6.6	SHAP Feature Importance	55
6.7	Window-Level versus Recording-Level Metrics	57
7	Analysis	60
7.1	Ablation Study	60
7.2	Performance Ceiling Analysis	61
7.3	Calibration Analysis	63
7.4	Error Analysis — Misclassified Recordings	64
7.5	Literature Benchmarking	67
7.6	Feature Interaction Analysis	68

8	Discussion	69
8.1	Performance in Clinical Context	69
8.2	Data-Level Bottleneck as the Binding Constraint	70
8.3	The Source-Confound Contribution	71
8.4	Clinical Interpretability	71
8.5	Limitations	72
8.6	Future Work	73
9	Conclusion	74
9.1	Summary of Work	74
9.2	Key Findings	75
9.3	Significance and Outlook	76
	Bibliography	78

List of Tables

2.1	<i>Comparative summary of machine learning methods for Brugada Syndrome classification (Part 1: Clinical targets and dataset sizes). The row for this thesis is highlighted as the primary contribution. . .</i>	11
2.1	<i>(Part 2: Architecture and performance). AUC values are as reported in the original works. Patient-disjoint CV denotes whether ECG recordings from the same patient were strictly confined to a single split. A dash (—) indicates not reported or not applicable.</i>	11
2.2	<i>Open gaps in the BrS classification literature and the corresponding contributions of this thesis. Each row maps a specific methodological or clinical gap, identified from the literature reviewed in this chapter, to the response implemented in the BrugadaECGNet pipeline.</i>	14
3.1	<i>Dataset composition. Native Fs: sampling rate of the raw source signal before preprocessing resampling. Final Fs: canonical sampling rate used for all downstream processing. Window estimates are based on the number of cardiac cycles detectable within each file’s signal duration: approximately 10–12 per 10-second Type-3 BrS recording, and approximately 3–4 per Non-BrS page ECG (typical digitized page duration: 2–4 seconds).</i>	16
3.2	<i>Per-fold patient-disjoint dataset statistics. NB = Non-BrS; T3 = Type-3 BrS. File counts reflect the number of ECG recordings (one prediction per file at recording level) assigned to each partition. No patient’s files appear in more than one partition within any fold. . .</i>	18
4.1	<i>PTB-XL RBBB-vs-NORM pretraining: epoch-by-epoch validation AUC. The encoder is trained on 48,444 training windows (3,369 NORM + 1,123 RBBB records) and evaluated on 23,082 validation windows (1,871 NORM + 285 RBBB records). The checkpoint saved at epoch 12 (val AUC = 0.9919) is used for weight transfer to BrugadaECGNet. Early stopping triggers at epoch 22 (patience = 10). . .</i>	24

4.2	Transfer learning ablation: Random initialisation, $D_{\text{feat}} = 1$ placeholder vs. PTB-XL pretrained encoder, same $D_{\text{feat}} = 1$ placeholder. All metrics are window-level, 5-fold patient-disjoint CV. $\Delta(\sigma)$ is the gain expressed in units of the standard deviation. Both changes satisfy the $\geq 1\sigma$ acceptance criterion.	25
4.3	BrugadaECGNet feature vector: all 44 features. NaN% is from V1 fold-0 train set (r' -dependent features); J-point-anchored features have 0% NaN. Clinical source column indicates the primary reference informing each feature. Features appearing in the GradientExplainer SHAP top-15 (fold-3) are marked †; the strongest individual binary discriminator by separation score is marked *.	27
4.4	Feature separation diagnostics on fold-0 training set, selected features. Separation = $ \text{median}_{\text{Type-3}} - \text{median}_{\text{Non-BrS}} $ normalised by the pooled interquartile range. Features near zero (Corrado index, ST curvature V1, fQRS count V1) were flagged as SHAP-prune candidates based on the GradientExplainer audit and their near-zero separation scores, but were retained after the pruning experiment failed to improve performance (see text). Note that ST curvature sign V2 and fQRS count V2 appear in the GradientExplainer top-15 (ranks 15 and 10 respectively), confirming that V1 and V2 variants carry different information.	28
4.5	Stream C addition results (Plan 16 Step 5) versus clean baseline (Plan 16 Step 0.5). All metrics are 5-fold patient-disjoint CV. RecT3F1 at Youden-J regresses due to threshold-transfer artefact (see text); RecT3F1 at T3-opt threshold recovers to 0.761 ± 0.043	32
4.6	Stream B widening + deeper fusion versus locked baseline (T3-opt threshold throughout). The formal $\geq 1\sigma$ criterion is not met on any primary metric, but no primary metric regresses and recording-level variance is nearly halved.	33
4.7	Development trajectory and ablation summary. Each row is a locked milestone. WinAUC and RecAUC are threshold-independent; RecAcc, RecNBF1, and RecT3F1 use the Youden-J threshold for pre-Phase 5 milestones (T3-opt not yet calibrated) and the T3-opt threshold from Phase 5 onwards. Dashes indicate metrics not yet available at that development stage. † Youden-J threshold. ‡ T3-opt threshold.	34
5.1	Input tensor channel assignment. All raw channels are mean-subtracted and divided by the standard deviation of the first 30 pre-QRS samples (per-channel baseline normalisation). Derivative channels are additionally normalised to unit standard deviation per channel. . . .	37

5.2	<i>BrugadaECGNet parameter count breakdown by component. Frozen components are transferred from the PTB-XL RBBB-vs-NORM pretrained encoder. Randomly initialised components include Stream B (44 features rather than 1 placeholder), Stream C (new), and the wider fusion head.</i>	41
6.1	<i>Five-fold patient-disjoint cross-validation aggregate performance. Window-level metrics use the per-fold window-level Youden-J threshold saved in <code>models/threshold_fold{N}.npy</code>. Recording-level T3-opt values are from <code>evaluation_results/cv_results.json</code>; recording-level Youden-J values are from the locked-baseline evaluation (recording-level Youden-J thresholds are calibrated on val recording ROC and not stored as a separate per-fold file). ROC-AUC is threshold-independent; the Youden-J AUC entry is omitted as it is identical to the T3-opt value. Recording-level T3-opt rows (bold) are the headline results.</i>	46
6.2	<i>Per-fold recording-level metrics at the T3-opt threshold. The T3-opt threshold is calibrated on the validation recording-level probabilities of each fold independently and recomputed at inference time; it is not stored as a separate per-fold file. RecAUC is threshold-independent. All values from <code>evaluation_results/cv_results.json</code>.</i>	47
6.3	<i>Attention weight entropy $H(\alpha)$ by class, aggregated across five folds and all test windows. The uniform ceiling $\log(53) = 3.970$ nats represents maximally diffuse attention. Both classes operate within 0.001–0.002 nats of the ceiling, confirming near-uniform attention behaviour.</i>	53
6.4	<i>Window-level versus recording-level metric comparison. Window-level uses the per-fold window-level Youden-J threshold; recording-level uses the T3-opt threshold. The Δ column shows the change from window to recording level. The AUC comparison (+0.018) is threshold-independent and is the most informative measure of aggregation benefit; the Accuracy and F1 deltas conflate the aggregation effect with the threshold change and should be interpreted with caution.</i>	58

7.1	<i>Ablation study: development trajectory from the two-stream Type-3 BrS baseline to the final three-stream locked architecture. Rows 1–2 are evaluated at the Youden-J threshold (T3-opt not yet implemented); rows 3–4 are evaluated at the T3-opt threshold. The Reverted row summarises four independently tested hypotheses (label smoothing; wider Stream B with reduced fusion bottleneck; learnable attention temperature; class-conditional pos_weight loss reweighting) that each produced a net metric regression or subthreshold gain and were reverted to the prior baseline. Δ columns are relative to the preceding kept row.</i>	61
7.2	<i>Per-fold recording-level error analysis at the T3-opt threshold . . .</i>	65
7.3	<i>Comparison of this thesis to published literature. AUC values are directly comparable only when the same ECG type and evaluation protocol are used; differences in BrS subtype, dataset, and validation methodology limit numerical comparison. N/A: not applicable; not stated: not reported in the source publication.</i>	67

List of Figures

2.1	Schematic comparison of the three Brugada Syndrome ECG pattern types in lead V1. (A) Type-1 (coved): J-point elevation ≥ 2 mm, convex descending ST segment, negative T-wave. (B) Type-2 (saddleback): J-wave elevation ≥ 2 mm, concave ST segment remaining ≥ 1 mm above baseline, positive T-wave. (C) Type-3 (saddleback/coved): ST elevation < 2 mm, most morphologically subtle. The α -angle, β -angle, and ST measurements at J, J+40 ms, and J+80 ms are indicated.	7
2.2	Standard 12-lead ECG electrode configuration showing the spatial organisation of the 16-channel BrugadaECGNet input. Leads V1, V2, V3, and V6 (Stream A, precordial encoder) sample the right ventricular outflow tract and lateral chest. Leads II, III, aVF, and aVR (Stream C, inferior encoder) sample inferior left ventricular and superior cardiac vectors, capturing the inferior ST-depression and aVR patterns associated with BrS risk stratification [11, 14]. . .	8
3.1	Source-confound audit: Pre-normalisation	19
3.2	Source-confound audit: Post-normalisation	20
3.3	Per-fold recording-level accuracy before and after.	20
4.1	SHAP feature importance bar chart	29
5.1	Input tensor construction pipeline.	38
5.2	BrugadaECGNet architecture diagram.	41
5.3	Training curves for a representative fold	43

6.1	Recording-level ROC curve from the production 80/10/10 patient-disjoint split.x-axis: False Positive Rate (1− Specificity); y-axis: True Positive Rate (Sensitivity). RecAUC = 0.963. The T3-opt operating point ($\tau = 0.297$, tuned on the production validation split) is marked. This figure depicts the production split and is presented as complementary to, not a replacement for, the five-fold CV aggregate RecAUC 0.939 ± 0.032 in Table 6.1.	48
6.2	Per-fold cross-validation recording-level ROC curves	49
6.3	Per-fold cross-validation window-level ROC curves	50
6.4	Window-level confusion matrices for the five cross-validation folds .	51
6.5	Recording-level confusion matrix from the production 80/10/10 patient-disjoint split	51
6.6	T3-opt vs. Youden-J threshold comparison (recording level)	52
6.7	Attention weight entropy	54
6.8	Representative per-window attention weight profiles	55
6.9	SHAP feature importance: top-15 features by mean SHAP value, computed using a GradientExplainer (<code>src/shap_feature_pruning.py</code>) on the fold-3 checkpoint (200 test windows, 100 background samples from the training set)	57
6.10	Window-level versus recording-level metric comparison	59
7.1	RecT3F1 sensitivity analysis	62
7.2	Calibration curves (reliability diagrams) for recording-level probabilities	64
7.3	SHAP beeswarm plot for the top-10 features by mean SHAP (GradientExplainer, fold-3 checkpoint, 200 test windows, 100 background samples). Each point represents one window; the x-axis is the SHAP value (positive $\rightarrow$ higher Type-3 BrS probability, negative $\rightarrow$ lower)	68

List of Acronyms

ADC

Analogue-to-Digital Converter

AUC

Area Under the ROC Curve

AUPRC

Area Under the Precision-Recall Curve

BCE

Binary Cross-Entropy

BCELoss

Binary Cross-Entropy Loss

BN

Batch Normalisation

BrS

Brugada Syndrome

CNN

Convolutional Neural Network

CRBBB

Complete Right Bundle Branch Block

CV

Cross-Validation

DL

Deep Learning

DWT

Discrete Wavelet Transform

ECG

Electrocardiogram

ECGD

ECG Digitizer (software)

FFN

Feed-Forward Network (within Transformer)

fQRS

Fragmented QRS

GE

General Electric

iRBBB

Incomplete Right Bundle Branch Block

LDA

Linear Discriminant Analysis

MAE

Major Arrhythmic Event

ML

Machine Learning

MLP

Multi-Layer Perceptron

Non-BrS

Non-Brugada Syndrome (negative class)

RBBB

Right Bundle Branch Block

RecAcc

Recording-level Accuracy (at T3-opt threshold)

RecAUC

Recording-level Area Under the ROC Curve

RecNBF1

Recording-level Non-BrS F1-score (at T3-opt threshold)

RecT3F1

Recording-level Type-3 BrS F1-score (at T3-opt threshold)

RF

Random Forest

ROC

Receiver Operating Characteristic

RVOT

Right Ventricular Outflow Tract

SCN5A

Sodium Channel Gene (Nav1.5)

SCD

Sudden Cardiac Death

SCP-ECG

Standardised Communications Protocol for ECG (PTB-XL annotation scheme)

SHAP

SHapley Additive exPlanations

SMOTE

Synthetic Minority Over-sampling Technique

ST

ST Segment (between QRS complex and T-wave)

T3-opt

Type-3-F1-optimising threshold (primary operating point)

TTA

Test-Time Augmentation

Type-3 BrS

Brugada Syndrome Type 3

WinAcc

Window-level Accuracy

WinAUC

Window-level Area Under the ROC Curve

WinNBF1

Window-level Non-BrS F1-score

Youden-J

Youden's J Statistic Threshold ($\text{argmax TPR} - \text{FPR}$)

Chapter 1

Introduction

1.1 Background and Clinical Motivation

Brugada Syndrome (BrS) is a rare but potentially fatal inherited cardiac channelopathy first described by Brugada and Brugada in 1992, who identified eight patients presenting with a history of aborted sudden cardiac death (SCD) and a distinctive electrocardiographic (ECG) pattern characterised by right bundle branch block morphology with persistent ST-segment elevation in the right precordial leads V1 to V3, in the absence of identifiable structural heart disease [1]. In the three decades since its original description, BrS has been recognised as one of the leading causes of sudden arrhythmic death in individuals with structurally normal hearts, accounting for an estimated 4–12% of all sudden death cases and approximately 20% of SCD events in patients under 50 years of age with no prior cardiac history [2].

The genetic basis of BrS is heterogeneous: loss-of-function mutations in the **SCN5A** gene account for approximately 20–30% of Western cases, with more than 20 additional implicated genes [2, 3]. BrS exhibits a pronounced male predominance (approximately 91.3% of symptomatic cases) and is most prevalent in Thailand and Southeast Asia, with global estimates of 1 in 2,000–5,000 individuals [2].

The electrocardiographic presentation of BrS has been formally classified into three pattern types, as codified by the 2013 HRS/EHRA/APHRS expert consensus statement on inherited primary arrhythmia syndromes [3] and subsequently refined in the 2016 expert consensus report on J-wave syndromes [4]. The **Type-1** (coved) pattern is characterised by a ≥ 2 mm J-point elevation in V1 or V2, a coved ST segment descending to an inverted T-wave, and is the sole ECG pattern pathognomonic for BrS; its spontaneous or drug-elicited presence is the primary diagnostic criterion [4]. The **Type-2** (saddleback) pattern is defined by a ≥ 2 mm J-wave with a concave ST segment remaining ≥ 1 mm above the isoelectric

line and a positive or biphasic T-wave [3]. The **Type-3** pattern encompasses either morphology in V1–V2 with ST elevation less than 2 mm, and is the most morphologically subtle and clinically ambiguous BrS manifestation — and the focus of this thesis [3].

The Type-3 pattern can spontaneously convert to a diagnostic Type-1 pattern during febrile illness, drug challenge, or autonomic fluctuation [3, 4]. Individuals with a Type-3 pattern are not classified as definitively BrS-positive without additional provocation, yet may harbour the same channelopathy and arrhythmic risk. Automated identification of Type-3 ECGs for referral to drug-challenge testing and genetic counselling is therefore a clinically meaningful and underserved problem, whose importance grows as risk stratification frameworks increasingly incorporate ECG biomarkers beyond the Type-1 pattern [5].

1.2 Limitations of Current ECG Diagnosis

Despite the clinical importance of accurate Type-3 BrS identification, the current diagnostic workflow relies predominantly on visual ECG interpretation by trained cardiologists. Published assessments of expert cardiologist performance on BrS classification tasks report AUC values in the range of 0.72–0.88, implying a substantial rate of misclassification even among specialist readers. This is a reflection not merely of individual skill variation but of the inherent morphological ambiguity of the Type-3 pattern itself, which overlaps visually with a range of benign and pathological conditions including early repolarisation syndrome, incomplete right bundle branch block, right ventricular hypertrophy, pericarditis, and the athlete’s heart adaptation. All of these conditions can produce right precordial ST elevation or J-wave deformations that are difficult to distinguish from Type-3 BrS without pharmacological provocation.

The ambiguity of the Type-3 morphology has been quantified in the electrocardiographic literature. Ohkubo et al. [6] demonstrated that the terminal QRS angle which is a morphological parameter measured in the frontal plane of the QRS complex in lead V1 is $38 \pm 24^\circ$ in Type-3 BrS and $38 \pm 14^\circ$ in Type-2 BrS, values that are statistically indistinguishable. This near-identical angular geometry confirms that the two patterns share fundamental morphological structure, and that Type-3 is not simply a lower-amplitude version of Type-2 that could be trivially detected by amplitude thresholding alone. In combination with the lower absolute ST elevation relative to the Type-1 pattern, this morphological similarity renders Type-3 discrimination particularly unreliable at the expert-reader level.

Critically, this limitation is structural. Current clinical guidelines do not provide definitive management recommendations for patients presenting with a Type-3 pattern in isolation, without a history of syncope, a documented arrhythmic event,

or a family history of BrS [3, 4]. The absence of an automated screening tool capable of reliably flagging Type-3 ECGs for further clinical evaluation means that a portion of at-risk individuals may remain unidentified until a sentinel arrhythmic event occurs. The development of a principled machine learning pipeline specifically designed to discriminate Type-3 BrS from Non-BrS ECGs therefore represents a direct response to an unmet clinical need.

1.3 Machine Learning for Brugada Syndrome: The Gap

The application of machine learning (ML) to automated ECG interpretation has expanded rapidly, and BrS has attracted increasing attention as a classification problem. Published work spanning CNN-based classifiers [7, 8] and transfer-learning pipelines [9] establishes a benchmark AUC range of 0.93–0.97, with PTB-XL pretraining identified as a consistently critical performance driver.

However, a fundamental gap persists: virtually all published ML pipelines target the Type-1 coved pattern, the most visually distinct BrS morphology. When Type-2 and Type-3 patterns are included they are typically subsumed into a single positive class, or addressed in a three-class formulation without patient-level data isolation [7]. The absence of patient-disjoint cross-validation is a pervasive limitation, inflating performance estimates by exploiting inter-window correlations across the same patient’s recordings. No published pipeline has been specifically designed and validated for the binary task of classifying Type-3 BrS against Non-BrS ECGs. This thesis directly addresses all of these gaps.

1.4 Contributions of This Work

This work presents five original contributions to the problem of automated Type-3 BrS classification from standard 12-lead ECGs:

1. **Dataset construction and curation.** A 237-file ECG dataset combining 174 digitised paper Non-BrS recordings (ECGD software) with 63 GE Sapphire DCAR XML Type-3 BrS recordings from a clinical centre in Torino, Italy. A patient manifest of 29 multi-file groups enables strict patient-disjoint evaluation.

2. **Source-confound identification and correction.** A three-diagnostic audit identified a 40-fold pre-QRS baseline variance difference between source populations. Per-channel baseline normalisation reduced this to $1.84\times$, with z-score ablation confirming the residual fingerprint contributes <0.005 to RecAUC. This method is transferable to any multi-source ECG study.

3. **A 44-dimensional clinically grounded feature vector.** Incorporating the de Luna α/β -angle descriptors [10], Corrado index, Crea inferior ST-depression

criterion [11], Babai-Bigi aVR sign, r-J interval V1 [12], and others. All 44 features were retained after pruning experiments confirmed low-SHAP features contribute through non-linear MLP interactions.

4. BrugadaECGNet: three-stream architecture with transfer learning.

Combining a PTB-XL pretrained CNN-Transformer precordial encoder (Stream A), a 44-feature MLP (Stream B), and an inferior-lead CNN (Stream C), fused with a T3-opt calibrated threshold. Across five patient-disjoint CV folds: RecAUC 0.939 ± 0.032 , RecAcc 0.872 ± 0.018 , RecT3F1 0.766 ± 0.034 — competitive with the published ML benchmark (0.93–0.97) and exceeding documented cardiologist AUC (0.72–0.88).

5. Empirical characterisation of the data-level performance ceiling.

Three post-baseline interventions (TTA, multi-seed ensemble, feature expansion) each failed the $\geq 1\sigma$ acceptance criterion, establishing that the RecT3F1 ceiling at ≈ 0.766 is a data-level constraint imposed by ≈ 15 –20 unique Type-3 patients, not an architectural limitation.

1.5 Thesis Outline

Chapter 2 provides clinical background and a structured literature review. Chapter 3 describes dataset construction, patient-disjoint cross-validation design, and the source-confound audit. Chapter 4 narrates the iterative pipeline development across six phases. Chapter 5 presents the final BrugadaECGNet architecture and training protocol. Chapter 6 reports cross-validation performance results. Chapter 7 provides ablation, ceiling, and error analyses. Chapter 8 discusses clinical implications, limitations, and future directions. Chapter 9 summarises the principal findings.

Chapter 2

Background and Literature Review

2.1 Brugada Syndrome — Clinical and Genetic Overview

Loss-of-function mutations in **SCN5A** account for 20–30% of definitive BrS diagnoses in Western cohorts; more than 20 additional genes have been implicated [2, 3]. BrS exhibits incomplete penetrance and variable expressivity: the ECG pattern may be concealed at rest and emerge only during fever, autonomic shifts, or pharmacological challenge. The arrhythmogenic mechanism is transmural dispersion of repolarisation in the RVOT epicardium, capable of triggering polymorphic ventricular tachycardia and ventricular fibrillation [2].

Risk stratification remains contested. The most established markers are spontaneous Type-1 pattern, symptomatic arrhythmia history, and male sex [3, 4]. Emerging frameworks add ECG biomarkers: Kan et al. [5] identified S-wave amplitude in Lead I (90.6% sensitivity), inferior early repolarisation, and PR prolongation; Ishida et al. [12] identified the r-J interval in V1 as the top SHAP feature in a Japanese BrS cohort. These biomarkers directly informed the 44-dimensional feature vector of this thesis.

2.2 ECG Morphology of Brugada Patterns

BrS diagnosis centres on leads V1 and V2, which overlie the RVOT. The key anatomical landmarks are the **J-point** (junction between QRS and ST segment), an accessory **r'-wave** deflection representing delayed RVOT depolarisation (present in saddleback patterns), and the **ST segment** whose amplitude, curvature, and

descent rate are the primary discriminating features. T-wave polarity in V1 is negative in Type-1 and positive or biphasic in Type-2/Type-3.

The three BrS ECG pattern types are defined by precise criteria as follows. The **Type-1 (coved) pattern** is characterised by a J-point elevation ≥ 2 mm in V1 or V2, a convex (coved) ST segment that descends monotonically from the J-point, and a negative T-wave; the overall morphology resembles an inverted triangle, with the apex at the J-point and the descending limb forming the coved ST segment. The Type-1 pattern is the only pattern considered pathognomonic for BrS under current diagnostic standards [4]. The **Type-2 (saddleback) pattern** is characterised by a J-wave elevation ≥ 2 mm producing a saddle-shaped ST segment that remains ≥ 1 mm above the isoelectric line before returning to a positive or biphasic T-wave. The r'-wave peak marks the high point of the initial elevation, and the ST segment has a concave or horizontal profile between the r'-peak and the T-wave onset. The **Type-3 pattern** is morphologically the most ambiguous: it encompasses either a saddleback or coved appearance in V1–V2 with ST elevation strictly less than 2 mm. Its definition is one of exclusion, specifically presenting a right precordial morphological abnormality in the BrS family that does not meet the amplitude criteria for Type-1 or Type-2 [3]. The morphological proximity of Type-3 to Type-2 has been quantified by Ohkubo et al. [6], who demonstrated a terminal QRS angle of $38 \pm 24^\circ$ in Type-3 and $38 \pm 14^\circ$ in Type-2 values that are statistically indistinguishable, motivating a feature-rich, multi-criterion approach to automated classification rather than reliance on any single morphological discriminator.

Several quantitative criteria capture morphological features beyond amplitude thresholds. The **de Luna α/β -angle system** [10] measures the ascending (α) and descending (β) slopes of the r'-wave triangle; a wider α characterises the saddleback and a steeper β the coved pattern (feature indices 0–3 in V1, 9–12 in V2). The **Corrado index** [13] is the ratio $ST(J)/ST(J+80ms)$: values >1 indicate coved morphology (BrS) while values ≤ 1 indicate benign early repolarisation (feature indices 6 and 15). The **Crea criterion** [11] flags ST depression ≥ 0.1 mV for ≥ 80 ms in leads II, III, and aVF, present in $\approx 47\%$ of Type-1 BrS patients and hypothesised to reflect global repolarisation dispersion (feature block 4, indices 29–40). The **Babai-Bigi aVR sign** [14] is defined as R-wave amplitude ≥ 0.3 mV or R/Q ratio ≥ 0.75 in aVR, associated with arrhythmic events in 84% of positive versus 27% of negative patients (feature block 5, indices 41–43).

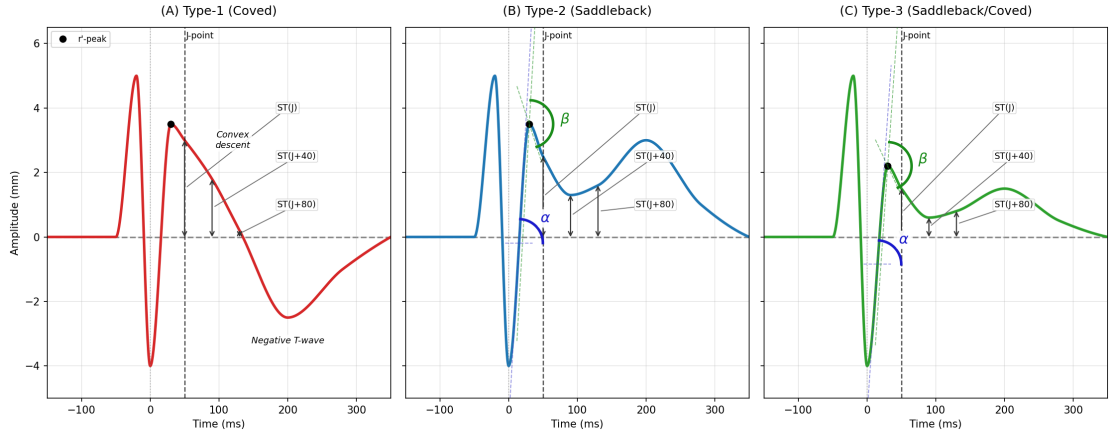


Figure 2.1: Schematic comparison of the three Brugada Syndrome ECG pattern types in lead V1. (A) Type-1 (coved): J-point elevation ≥ 2 mm, convex descending ST segment, negative T-wave. (B) Type-2 (saddleback): J-wave elevation ≥ 2 mm, concave ST segment remaining ≥ 1 mm above baseline, positive T-wave. (C) Type-3 (saddleback/coved): ST elevation < 2 mm, most morphologically subtle. The α -angle, β -angle, and ST measurements at J, J+40 ms, and J+80 ms are indicated.

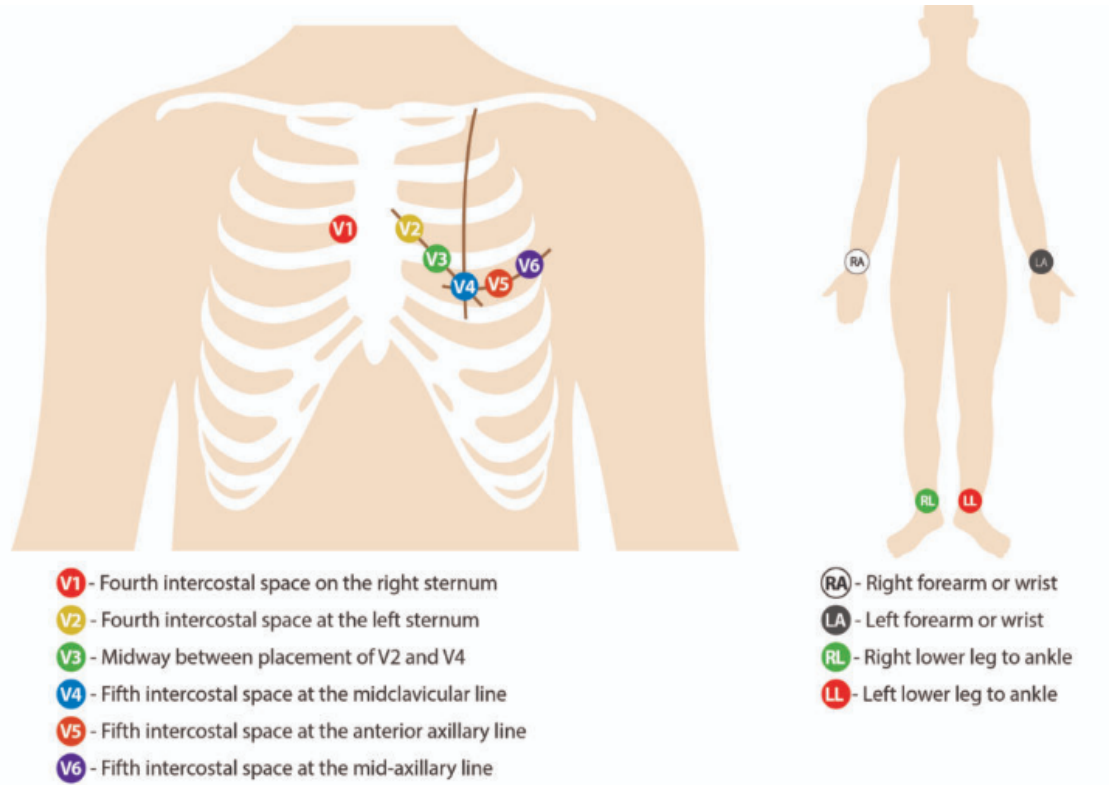


Figure 2.2: Standard 12-lead ECG electrode configuration showing the spatial organisation of the 16-channel BrugadaECGNet input. Leads V1, V2, V3, and V6 (Stream A, precordial encoder) sample the right ventricular outflow tract and lateral chest. Leads II, III, aVF, and aVR (Stream C, inferior encoder) sample inferior left ventricular and superior cardiac vectors, capturing the inferior ST-depression and aVR patterns associated with BrS risk stratification [11, 14].

2.3 Prior Machine Learning Work for Brugada Classification

The application of machine learning to automated BrS classification has expanded appreciably over the preceding five years, driven by increasing availability of digitised 12-lead ECG repositories and advances in deep learning architectures for time-series data. The literature spans CNN-based approaches on raw ECG signals, deep learning pipelines augmented by transfer learning, systematic reviews establishing performance benchmarks, and clinical feature-based risk stratification studies. A structured review is provided below, followed by a comparative summary in Table 2.1.

2.3.1 CNN-Based Approaches

Randazzo et al. [7] presented a three-class CNN-based BrS classifier trained on 349 ECG recordings. Signal preprocessing used an S-peak-centred 850 ms window (300 ms pre-S, 550 ms post-S) and Gaussian noise augmentation ($\sigma = 0.15$), achieving a macro F1-score of 0.93. Both the S-peak window and noise augmentation were adopted in the pipeline of this thesis. Notably, the non-Type-1 class conflated Type-2 and Type-3 patterns, and patient-disjoint cross-validation was not employed.

Melo et al. [8] applied deep learning to 1,154 ECG recordings including an external validation cohort, achieving an AUC of 0.934. Beat-averaging was used for noise reduction; an analogous multi-window probability aggregation was evaluated in this thesis and found to provide no measurable improvement under patient-disjoint CV. The work targets Type-1 BrS predominantly and does not employ patient-disjoint evaluation.

2.3.2 Deep Learning with Transfer Learning

Saleh et al. [9] achieved an AUC of 0.981 on 704 ECG recordings using supervised pretraining on the PTB-XL database (RBBB-versus-normal task). RBBB produces delayed right ventricular depolarisation in V1/V2, directly adjacent to the BrS morphological region, making it the closest available pretraining proxy. Saleh et al. reported a +3.2% AUC improvement from PTB-XL pretraining relative to random initialisation; SMOTE was found inconsistent. The present thesis adopts the same approach, confirmed by a controlled ablation in Chapter 4.

2.3.3 Benchmarks and the Patient-Disjoint Gap

Across the reviewed ML-for-BrS literature, reported AUC values span 0.93–0.97, with cardiologist reader performance estimated at 0.72–0.88. A pervasive methodological limitation is the absence of patient-disjoint cross-validation: most published studies use simple random splits, yielding optimistic performance estimates that may not generalise to unseen patients. No published study directly targets the binary Type-3 BrS versus Non-BrS classification task under a patient-disjoint evaluation protocol — the precise gap addressed by this thesis.

2.3.4 Risk Stratification and Clinical Feature Studies

Ishida et al. [12] applied gradient-boosted machine learning to a Japanese BrS cohort of 234 patients with documented follow-up, identifying 29 major arrhythmic events. Using SHAP analysis to rank feature contributions, the study identified the r-J interval in lead V1 as the top predictive feature for arrhythmic event risk. This finding was incorporated into the 44-feature vector of this thesis as a cross-lead

global feature (index 23), providing a clinically validated morphological signal that would not be captured by amplitude-only ST measurements. The r-J interval reflects the degree of QRS terminal prolongation in V1, which is physiologically linked to the degree of RVOT conduction delay and therefore to the magnitude of the Phase 2 re-entry substrate.

Kan et al. [5] conducted an updated review of risk stratification markers for BrS, confirming the prognostic value of S-wave amplitude in Lead I (sensitivity 90.6% for arrhythmic events), early repolarisation in the inferior leads, and PR interval prolongation. These findings support the inclusion of the Crea inferior ST-depression flags (feature block 4) and the Babai-Bigi aVR sign features (feature block 5) in the 44-dimensional feature vector which are features that are motivated not only by clinical prevalence data but by independent evidence of their arrhythmic risk relevance.

Ohkubo et al. [6] measured the terminal QRS angle in V1 for 46 BrS patients, finding values of $38 \pm 24^\circ$ for Type-3 and $38 \pm 14^\circ$ for Type-2, a statistically indistinguishable difference that quantitatively establishes the morphological proximity of the two patterns. This result motivates the use of a rich, multi-criterion feature vector combining amplitude, slope, angular, and temporal features rather than any single discriminator, and further contextualises why the data-level bottleneck, arising from the small number of unique Type-3 patients available, is the principal limitation on achievable RecT3F1 at the current dataset scale.

Table 2.1: *Comparative summary of machine learning methods for Brugada Syndrome classification (**Part 1**: Clinical targets and dataset sizes). The row for this thesis is highlighted as the primary contribution.*

Reference	Year	Target ECG type(s)	Dataset size
Randazzo et al. [7]	2025	Type-1, non-Type-1 BrS, Non-BrS (3-class)	349 recordings
Melo et al. [8]	2023	Type-1 BrS vs. Non-BrS	1,154 recordings
Saleh et al. [9]	2024	BrS vs. Non-BrS	704 recordings
Ishida et al. [12]	2025	MAE risk (BrS cohort only)	234 BrS patients
This thesis	2026	Type-3 BrS vs. Non-BrS (binary)	237 recordings

Table 2.1: (**Part 2**: Architecture and performance). *AUC values are as reported in the original works. Patient-disjoint CV denotes whether ECG recordings from the same patient were strictly confined to a single split. A dash (—) indicates not reported or not applicable.*

Reference	Architecture	AUC	Patient-disjoint CV	Type-3 specific
Randazzo et al. [7]	CNN, S-peak window, noise augmentation	0.93 (macro F1)	No	No
Melo et al. [8]	Deep CNN, beat-averaging	0.934	No	No
Saleh et al. [9]	CNN + PTB-XL transfer learning	0.981	No	No
Ishida et al. [12]	Gradient boosting + SHAP	—	No	No
This thesis	BrugadaECGNet:CNN Transformer + MLP + Inferior CNN, PTB-XL transfer	0.939 $\pm$ 0.032 (RecAUC)	Yes (5-fold)	Yes

2.4 Attention Mechanisms and Transformer Architectures for ECG

The choice of a Transformer encoder as the core temporal sequence model in BrugadaECGNet is grounded in both the theoretical properties of the architecture and the empirical failure of the preceding BiLSTM-based approach during pipeline development.

The Transformer architecture, introduced by Vaswani et al. [15], replaces the sequential hidden-state propagation of recurrent models with a multi-head self-attention mechanism. Each position in the input sequence attends simultaneously to all other positions, computing weighted combinations of value vectors where the weights are determined by scaled dot-product similarity between query and

key representations. This global receptive field, combined with a position-wise feed-forward network applied identically at each timestep, enables the Transformer to model complex temporal dependencies without the sequential bottleneck inherent in recurrent architectures. Sinusoidal positional encodings are added to preserve temporal order information.

In the early development phases of this pipeline, a BiLSTM encoder [16] paired with an additive attention layer [17] was employed. A critical empirical failure motivated its replacement: attention weights consistently converged to the uniform ceiling ($H(\alpha) = \log(53) = 3.970$ nats) regardless of architectural variant. The root cause is information diffusion in BiLSTM hidden states — all positions become near-identical representations, leaving no score differences for attention to exploit. The Transformer’s self-attention avoids this, as each timestep retains a distinct representation (detailed in Chapter 4).

In the final BrugadaECGNet architecture, a 2-layer Transformer encoder processes the temporal sequence produced by the convolutional blocks. Following the Transformer, a custom hierarchical attention layer computes a single context vector $\mathbf{c} \in \mathbb{R}^{64}$ via a trainable context query, providing a domain-adapted temporal aggregation that distills the encoder output into a fixed-length feature vector for downstream fusion.

2.5 Transfer Learning and Domain Adaptation in ECG

Transfer learning [18] refers to the practice of leveraging knowledge acquired from a source task to improve performance on a related target task, typically by initialising the target model with weights pretrained on the source. In the context of deep neural networks, the transferability of features is a function of network depth. Early layers learn generic, broadly applicable representations in the ECG domain. These correspond to waveform morphology templates such as R-wave shapes, S-wave deflections, and ST-segment profiles, while deeper layers encode progressively more task-specific representations. Yosinski et al. [19] demonstrated empirically that freezing early transferred layers and fine-tuning only later layers is optimal when the source and target domains are related but not identical, as freezing preserves the generic representations while permitting task adaptation in the upper network.

The PTB-XL database [20] comprises 71,000 12-lead recordings from 18,885 patients at 500 Hz, providing a pretraining corpus several orders of magnitude larger than the 237-recording BrS dataset [21]. The pretraining task selected is RBBB-versus-normal binary classification. RBBB produces a broad, notched terminal R'-wave in V1/V2 arising from the same RVOT anatomical region as the Brugada ST abnormalities, making it the clinically closest available pretraining

proxy. Saleh et al. reported a +3.2% AUC improvement from PTB-XL pretraining relative to random initialisation; the present thesis adopts the same approach and confirms its contribution via a controlled ablation (Chapter 4). Negative transfer [18] is mitigated by selectively freezing the early convolutional filters and first Transformer layer after transfer, while upper task-specific layers fine-tune freely.

2.6 Multi-Source ECG Dataset Challenges and Confound Mitigation

Assembling training datasets from multiple clinical centres or acquisition systems is a pervasive challenge in clinical machine learning. Systematic differences between data sources, arising from hardware characteristics, signal processing pipelines, and recording protocols, can act as spurious discriminative features that a model exploits in lieu of clinically meaningful morphological patterns [22]. This phenomenon, known as dataset shift or source confounding, is particularly acute in ECG analysis. Signal amplitude, baseline noise floor, and frequency response are all sensitive functions of the specific acquisition hardware and software chain. When each class in a classification task is predominantly or exclusively associated with a different acquisition source, the classifier may learn to discriminate acquisition fingerprints rather than pathological morphology which is a failure mode that produces inflated in-sample performance but poor clinical generalisation [22].

In this thesis, the Type-3 BrS class (GE MAC2000 hospital-grade digital recordings) and the Non-BrS class (ECGD paper-ECG digitisations) originate from structurally different acquisition pipelines, giving rise to a systematic noise-floor disparity between classes. A three-diagnostic audit confirmed that the fingerprint was present but secondary to morphology, and per-channel baseline normalisation was developed as a principled, anatomy-preserving correction applied exclusively to the ECG tensor, leaving the millivolt-scale features of Stream B unaltered. The full audit methodology and results are presented in Chapter 3.

The open gaps identified across these sections, and the specific contributions of this thesis in response to each, are summarised in Table 2.2.

Table 2.2: *Open gaps in the BrS classification literature and the corresponding contributions of this thesis. Each row maps a specific methodological or clinical gap, identified from the literature reviewed in this chapter, to the response implemented in the BrugadaECGNet pipeline.*

Gap identified in literature	Primary source(s)	Response in this thesis
No ML pipeline specifically targets binary Type-3 BrS vs. Non-BrS	[7]	Pipeline designed specifically for Type-3 BrS vs. Non-BrS binary classification from the outset
Patient-disjoint CV rarely used; performance estimates are optimistic	Prior literature	5-fold patient-disjoint CV; all recordings from the same patient confined to one fold across train/val/test
BiLSTM attention entropy saturates at uniform ceiling in ECG encoding	Development (this thesis)	Replaced with 2-layer Transformer encoder; uniform entropy failure documented and resolved
PTB-XL transfer learning not applied to Type-3 BrS classification	[9]	RBBB-vs-normal PTB-XL pretraining; $+1.37\sigma$ AUC gain confirmed by controlled ablation
Multi-source amplitude fingerprints unaddressed in small BrS datasets	[22]	Three-diagnostic source-confound audit; per-channel baseline normalisation reducing fingerprint $40\times$ to $1.84\times$
Clinical criteria (Crea, Babai-Bigi, r-J V1) not combined in a single ML feature vector	[11, 14, 12]	44-dimensional feature vector incorporating all five clinical feature blocks in Stream B
Data-level performance ceiling for Type-3 BrS not empirically characterised	—	Systematic three-intervention study (TTA, multi-seed ensemble, feature expansion) confirming ceiling is data-level, not architectural

The following chapters present the implementation of each of these contributions in full detail, beginning with the dataset construction and curation process in Chapter 3.

Chapter 3

Dataset and Data Preparation

3.1 Digitized Paper ECGs

The Non-BrS class was assembled from approximately 1,000 standard paper 12-lead resting ECGs from a clinical archive covering diverse cardiac and non-cardiac diagnoses. A clinically heterogeneous negative class is deliberate: it requires the classifier to distinguish Type-3 BrS from the full morphological spectrum encountered in practice. Each recording was digitized using ECGD software, which outputs multi-page Excel files (.xlsx) with lead amplitudes in millivolts. Native sampling rates average ≈ 588 Hz and are resampled to 500 Hz (consistent with PTB-XL) before downstream processing. After clinical labelling and curation (Type-1 BrS files discarded), **174 Non-BrS files** were retained.

3.2 Type-3 BrS Class: GE Sapphire DCAR XML

The **63 Type-3 BrS recordings** were provided by a clinical centre in Torino, Italy, acquired on a GE MAC2000 ECG machine and stored in GE Sapphire DCAR XML format. Key specifications: 1,000 Hz sampling rate, signed 16-bit ADC integers at $4.88 \mu\text{V}/\text{count}$, applied conversion $\text{mV} = \text{ADC} \times 4.88/1000$, onboard 0.56 Hz/50 Hz/150 Hz filters, 25 mm/s, 10 mm/mV gain — matching Non-BrS class parameters. A Python conversion script exported signals to Excel; all subsequent pipeline steps (resampling to 500 Hz, S-peak detection, windowing, normalisation, feature extraction) are identical for both classes.

3.3 Clinical Labelling and Dataset Curation

Clinical labelling was performed by cross-referencing all digitized and converted ECG files against a structured diagnostic report (‘reports/dataset_report.txt’) covering the original ECG collection. The report enumerated 215 valid labelled instances corresponding to 165 distinct ECG case identifiers present in the physical archive. Each case was assigned one of three clinical labels: Type-1 BrS, Type-3 BrS, or Non-BrS. The label assignment was based on cardiologist interpretation documented in the clinical record associated with each ECG encounter.

The curation process involved three steps. First, each file in the digitized archive was matched to its corresponding entry in the diagnostic report using a normalised filename-to-case mapping, with manual correction applied to resolve cases where filename tokenisation produced ambiguous matches (e.g. the file ‘Capussotti marzo 2010.xlsx’ was assigned to patient group ‘Capussotti’ rather than the literal token ‘Capussotti marzo’). Second, files carrying a Type-1 BrS label were removed from the dataset entirely, as the thesis targets the clinically distinct and more diagnostically challenging binary discrimination between Type-3 BrS and Non-BrS. Inclusion of Type-1 patterns would trivially inflate classifier performance due to the visually unambiguous coved ST morphology of Type-1. Third, files with no matching label in the diagnostic report, arising from recording dates or patient identifiers absent from the report were excluded from the labelled dataset and not used for training or evaluation.

Following curation, the final labelled dataset comprises **174 Non-BrS files** and **63 Type-3 BrS files**, totalling **237 ECG recordings** from **83 unique patient groups**. The class composition, source format, approximate window yield, unique patient count, and sampling rates for each class are summarised in Table 3.1.

Table 3.1: *Dataset composition. Native Fs: sampling rate of the raw source signal before preprocessing resampling. Final Fs: canonical sampling rate used for all downstream processing. Window estimates are based on the number of cardiac cycles detectable within each file’s signal duration: approximately 10–12 per 10-second Type-3 BrS recording, and approximately 3–4 per Non-BrS page ECG (typical digitized page duration: 2–4 seconds).*

Class	Source format	Files	Approx. windows	Unique patients	Native Fs (Hz)	Final Fs (Hz)
Type-3 BrS	GE Sapphire DCAR XML (converted to XLSX)	63	~700	54	1,000	500
Non-BrS	ECGD paper digitization (multi-sheet XLSX)	174	~600	29	~588 (variable)	500
Total	—	237	~1,300	83	—	500

A notable structural asymmetry characterises the two classes: the 63 Type-3 BrS files originate from 54 distinct patients, with only 3 patients contributing more than one recording (predominantly single-visit recordings). By contrast, the 174 Non-BrS files originate from just 29 distinct patients, with 26 of those patients contributing multiple recordings spanning different clinical encounters and dates. This asymmetry arises from the nature of the two archival sources: the Type-3 BrS dataset was assembled as a targeted collection of Brugada patients from a single clinical centre, whereas the Non-BrS archive reflects a longitudinal clinical database in which individual patients had repeated ECG evaluations over multi-year follow-up periods. The implication for cross-validation design is discussed in Section 3.4.

3.4 Patient Grouping and Cross-Validation Design

The structural asymmetry described in Section 3.3, in which a small number of Non-BrS patients contribute a disproportionately large number of ECG files, makes patient-disjoint data splitting essential. Multiple files from the same patient share the same underlying cardiac physiology, body habitus, and recording environment. Their ECG morphologies are therefore substantially correlated. If files from the same patient are distributed across training and test sets, the model’s test performance is inflated by implicit memorisation of patient-specific features rather than generalisable class-discriminative patterns. A patient-disjoint split guarantees that all recordings from any single patient appear exclusively in one partition (train, validation, or test) across every fold.

Patient grouping was performed by applying regular expression rules to extract a patient identifier from each filename, handling eight distinct Non-BrS naming conventions. Three edge cases required manual correction. The resulting manifest maps all 237 filenames to one of 83 unique patient identifiers, with 29 multi-file patient groups confirmed.

The five-fold patient-disjoint cross-validation harness divides patient groups into five balanced chunks per class using a seeded greedy algorithm, ensuring no patient’s files span more than one partition. For fold i , chunk i is the test set, chunk $(i + 1) \bmod 5$ is validation, and the remaining three chunks form training — so every patient group appears exactly once in each role across the five folds. Per-fold file counts are reported in Table 3.2.

Table 3.2: *Per-fold patient-disjoint dataset statistics. NB = Non-BrS; T3 = Type-3 BrS. File counts reflect the number of ECG recordings (one prediction per file at recording level) assigned to each partition. No patient’s files appear in more than one partition within any fold.*

Fold	Train files (NB / T3)	Val files (NB / T3)	Test files (NB / T3)	Total test
0	140 (103 / 37)	48 (35 / 13)	49 (36 / 13)	49
1	142 (105 / 37)	47 (34 / 13)	48 (35 / 13)	48
2	144 (106 / 38)	46 (34 / 12)	47 (34 / 13)	47
3	144 (105 / 39)	47 (35 / 12)	46 (34 / 12)	46
4	141 (103 / 38)	49 (36 / 13)	47 (35 / 12)	47

The test partitions range from 46 to 49 files, with Type-3 BrS test recordings numbering 12–13 per fold. This small Type-3 test set size has a direct and quantifiable effect on recording-level metric resolution, with 12–13 Type-3 test recordings per fold, a single misclassification shifts the recording-level Type-3 F1-score (RecT3F1) by approximately 0.08. A granularity constraint is discussed in full in Chapter 6 when interpreting the performance ceiling.

The production model uses a single patient-disjoint 80/10/10 split rather than 5-fold CV. The same LPT chunk algorithm is applied with 10 balanced chunks per class (seed 42 for Non-BrS, seed 43 for Type-3 BrS). Chunks 0–7 form the training set (80%), chunk 8 the validation set (10%), and chunk 9 the test set (10%). The use of 10 chunks rather than 5 means the production split does not coincide exactly with any single CV fold, providing an independent evaluation on a held-out patient subset.

3.5 Source-Confound Audit

Because the Type-3 BrS class (hospital-grade GE MAC2000 digital recordings) and the Non-BrS class (paper-ECG ECGD digitisations) originate from structurally different acquisition pipelines, a source-confound audit was conducted prior to architecture development to quantify the risk of fingerprint exploitation.

Audit methodology. Three independent diagnostics were applied: (D1) Pearson correlation between pre-QRS baseline statistics and model logits; (D2) z-score ablation AUC drop; and (D3) within-class recording probability dispersion.

Pre-normalisation results (D1 fail). D1 confirmed fingerprint leakage: $|r| = 0.200$ in V1 and up to 0.331 in V2 ($p < 0.001$), caused by a 40-fold disparity

in V1 pre-QRS baseline variance (0.002761 mV^2 Non-BrS vs. 0.000070 mV^2 Type-3 BrS). D2 showed only $\Delta\text{AUC} = 0.016$ under z-scoring (D2 pass, fingerprint secondary to morphology), and D3 passed. The D1-fail/D2-pass/D3-pass verdict confirmed the fingerprint was present but secondary.

Normalisation fix. Per-channel baseline normalisation was applied: each of the 16 ECG channels is divided by the standard deviation of its first 30 pre-QRS samples, rescaling to noise-floor units. This is applied exclusively to the ECG tensor (Streams A and C); the hand-engineered features of Stream B retain the raw millivolt-scale signal required by clinical criteria.

Post-normalisation results (PASS). The baseline-variance ratio collapsed from $40\times$ to $1.84\times$. A repeat z-score ablation produced $\Delta\text{RecAUC} = +0.005$ (noise level), and all D1–D3 diagnostics passed. Recording-level accuracy improved by 0.030 and fold-to-fold standard deviation approximately halved; the fold 4 outlier (RecAcc 0.957) reverted to 0.872, confirming the outlier was fingerprint-driven. The source-confound correction, patient-disjoint evaluation, and dataset curation together form the empirical foundations for the model development presented in Chapters 4 and 5.

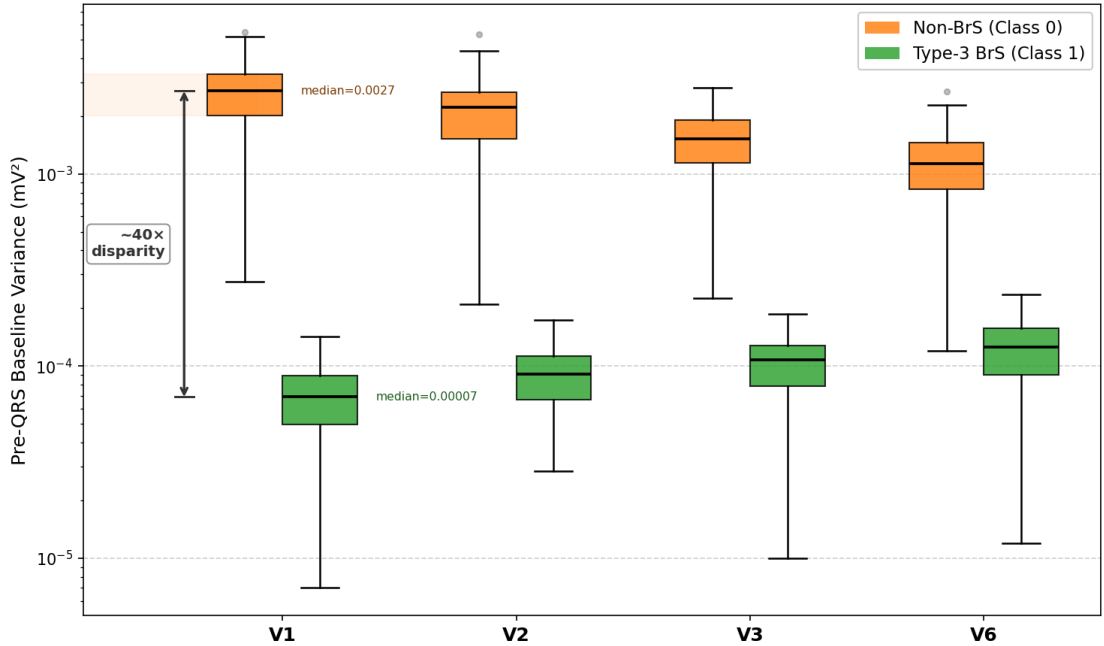


Figure 3.1: Source-confound audit: Pre-normalisation

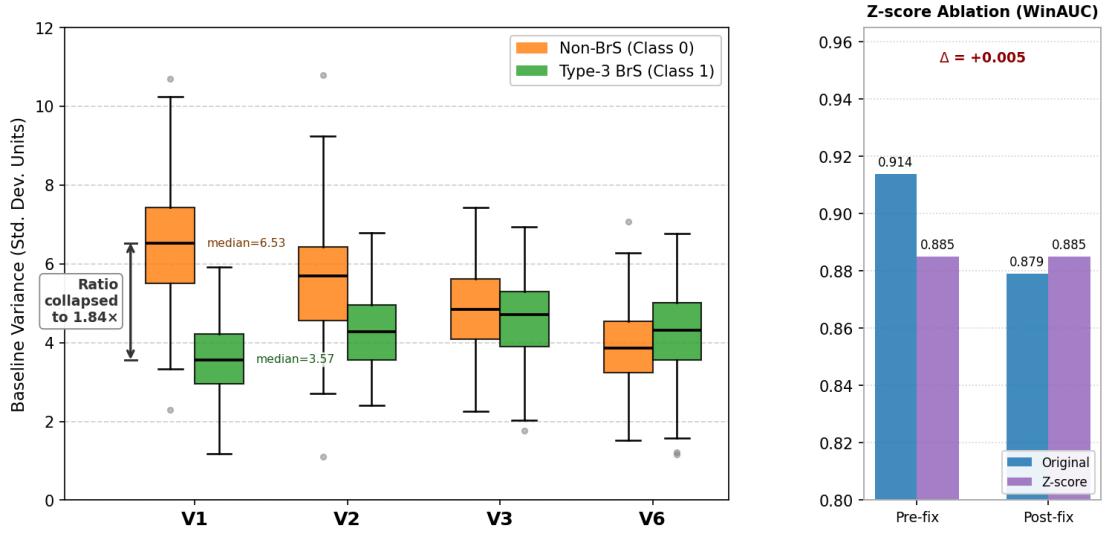


Figure 3.2: Source-confound audit: Post-normalisation

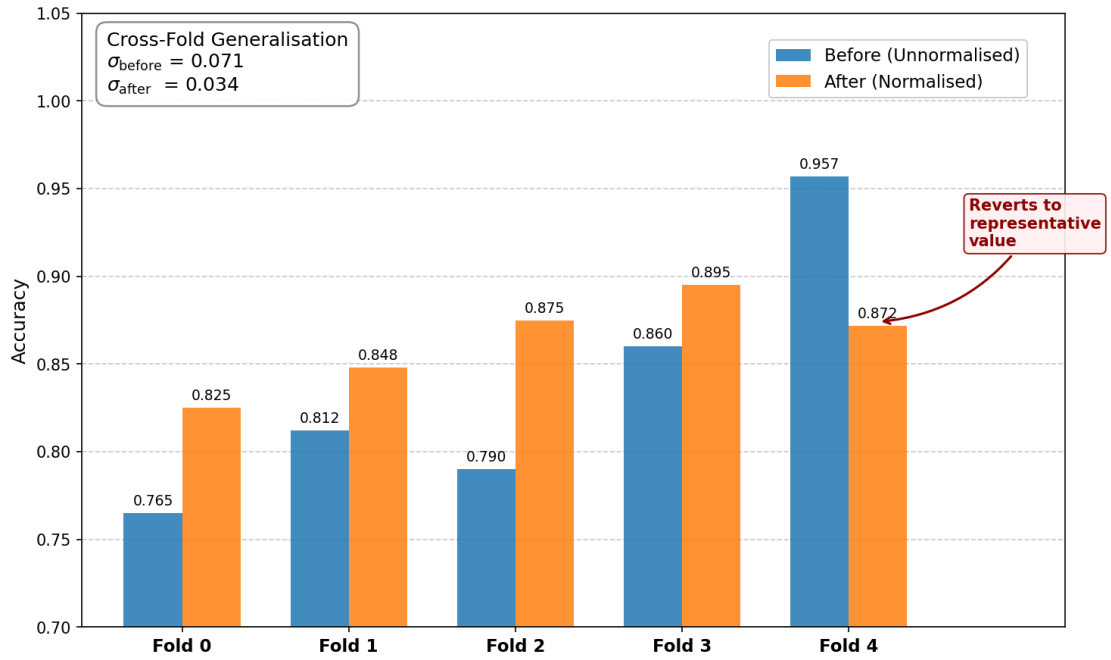


Figure 3.3: Per-fold recording-level accuracy before and after.

Chapter 4

Pipeline Development

4.1 Design Principles and Evaluation Criteria

The pipeline described in this thesis was developed through an iterative, experiment driven methodology. Each architectural intervention, feature addition, or data processing change was evaluated by running the full 5-fold patient-disjoint cross-validation and comparing the resulting metrics against the preceding baseline. A proposed change was retained only if it produced an improvement of at least one standard deviation ($\geq 1\sigma$) on the primary metric across folds, otherwise it was reverted and documented as a negative result. This acceptance criterion is presented here not as bureaucratic conservatism but as a methodological necessity dictated by the dataset size.

With only 237 ECG recordings distributed across five folds, each test partition contains approximately 10–13 Type-3 BrS recordings and 34–36 Non-BrS recordings. In a dataset of this scale, the assignment of a single patient group to one fold rather than another was driven by the seeded LPT chunk algorithm rather than by any property of the data which can shift recording-level metrics by as much as 0.05–0.08 in isolation. A change that produces, for example, a +0.03 lift in mean RecAcc across folds is well within the range of fold-assignment noise and provides no reliable evidence of generalisation. The $\geq 1\sigma$ gate operationalises the intuition that only changes whose benefit is large relative to the natural fold-to-fold variance are worth retaining in the final system. The rule was applied to window-level accuracy and Non-BrS F1-score in the early phases (prior to recording-level infrastructure) and to recording-level RecAcc, RecNBF1, and RecT3F1 in later phases.

The chapter that follows narrates the six development phases in sequence, presenting both the positive decisions retained in the final pipeline and the negative results that defined its boundaries.

4.2 Phase 1 — Preprocessing Infrastructure and Dataset Integration

Phase 1 established four core preprocessing components. **S-peak detection and windowing.** R-peaks are located via NeuroKit2 and S-peaks via DWT delineation, with an argmin fallback for robustness on noisy digitised recordings. An asymmetric 850 ms window (300 ms pre-S, 550 ms post-S = 425 samples at 500 Hz) is centred on each S-peak. The asymmetric split places the J-point and ST-segment morphology — the BrS discriminative region — in the central window zone; R-centring would shift this region toward the edge. This windowing matches the approach of Randazzo et al. [7].

Multi-format XLSX loading. Non-BrS multi-sheet workbooks and Type-3 BrS single-sheet workbooks are handled by two dedicated loading branches; missing lead columns are zero-filled.

ADC unit mismatch correction. An early development run yielded 0.98 validation accuracy — a red flag. The GE Sapphire XML files store raw ADC counts ($\pm 32,768$), not millivolts, while ECGD files are already in millivolts. The classifier was trivially separating classes by amplitude scale. After applying $\text{mV} = \text{ADC} \times 4.88/1000$ to all Type-3 BrS files, accuracy dropped to 0.72, confirming morphologically comparable signals.

Augmentation, patient manifest, and cross-validation harness. Gaussian noise with standard deviation $\sigma = 0.15$ is applied to the 16-channel input tensor during training only, following the augmentation protocol of Randazzo et al. [7]. The patient manifest and 5-fold patient-disjoint cross-validation harness (described in detail in Chapter 3) were also established in this phase, providing the unbiased evaluation infrastructure used for all subsequent experiments.

4.3 Phase 2 — Architecture Development: From BiLSTM to Transformer

The initial architecture was a three-channel BiLSTM (V1, V2, V3) with hierarchical attention. It failed persistently: attention entropy was pinned at $\ln(53) = 3.970$ nats (the uniform ceiling) in all configurations, indicating the model computed an unweighted mean of hidden states rather than attending to the ST-segment.

Seven architectural variants were attempted (Xavier initialisation, dimension reduction, entropy penalty, L2 variations, MaxPool resolution, weighted sampling, per-file window cap) — none escaped the ceiling. The failure is structural: BiLSTM hidden states diffuse information in both directions, making all 53 post-pooling representations near-identical, so pre-softmax attention logits are near-zero and cannot be focused by any downstream transformation.

Transition to Transformer encoder. A two-layer `TransformerEncoder` ($d_{\text{model}} = 64$, 4 heads, FFN=128, dropout=0.2) with sinusoidal positional encoding [15] replaced the BiLSTM. Transformer output representations are content-specific (each token attends to all others by query-key similarity), providing non-uniform pre-softmax inputs to the retained hierarchical attention layer, which pools the 53-timestep sequence to a 64-dim context vector.

8-channel precordial input. The input was extended from 3 raw leads (V1, V2, V3) to 8 channels: four raw precordial leads (V1, V2, V3, V6) plus their first-derivative channels, each unit-std normalised. Derivative channels directly encode the slope-based de Luna α/β -angle features [10]; V6 is required for the QRS mismatch feature.

4.4 Phase 3 — PTB-XL Transfer Learning

With only 237 recordings, the encoder is severely under-constrained for random-initialisation training. PTB-XL [20] pretraining on an RBBB-vs-NORM binary task was adopted as a mandatory foundation, following Saleh et al. [9] who demonstrated a +3.2% AUC benefit. The clinical rationale is domain proximity: RBBB produces broad, notched R'-waves in V1/V2 arising from the same RVOT anatomy as Brugada ST abnormalities.

Pretraining setup. PTB-XL stratified folds 1–8 for training, 9–10 for validation; NORM undersampled to $3\times$ RBBB count, yielding 48,444 training and 23,082 validation windows after the same 850 ms S-peak windowing pipeline. Best validation AUC of **0.9919** at epoch 12 (Table 4.1); early stopping triggered at epoch 22.

Table 4.1: *PTB-XL RBBB-vs-NORM pretraining: epoch-by-epoch validation AUC. The encoder is trained on 48,444 training windows (3,369 NORM + 1,123 RBBB records) and evaluated on 23,082 validation windows (1,871 NORM + 285 RBBB records). The checkpoint saved at epoch 12 (val AUC = 0.9919) is used for weight transfer to BrugadaECGNet. Early stopping triggers at epoch 22 (patience = 10).*

Epoch	Train loss	Val AUC	Saved
1	0.1526	0.9899	✓
2	0.1074	0.9906	✓
3	0.0972	0.9908	✓
4	0.0949	0.9908	—
5	0.0893	0.9905	—
6	0.0855	0.9907	—
7	0.0825	0.9894	—
8	0.0846	0.9912	✓
9	0.0797	0.9917	✓
10	0.0758	0.9918	✓
11	0.0757	0.9915	—
12	0.0729	0.9919	✓ best
13–22	0.0689–0.0512	0.9885–0.9912	—

Weight transfer and layer freezing. Seventy-two parameter keys are transferred to BrugadaECGNet’s Stream A using `load_state_dict(strict=False)`, covering both convolutional blocks, the linear projection, positional encoding, both Transformer layers, and the custom attention layer. Of these, 28 tensors are subsequently frozen (convolutional blocks 1–2 and Transformer layer 0), preserving the low-level and contextual morphology features learned on RBBB, while upper layers adapt to the BrS task.

Transfer learning ablation. The causal contribution of PTB-XL pretraining was quantified by a controlled ablation comparing randomly initialised encoder, $D_{\text{feat}} = 1$ placeholder feature against PTB-XL pretrained encoder, same $D_{\text{feat}} = 1$ placeholder. Results are reported in Table 4.2. The pretrained encoder produced gains of $+1.42\sigma$ in window-level accuracy ($0.527 \pm 0.083 \rightarrow 0.645 \pm 0.047$), $+1.03\sigma$ in Non-BrS F1-score ($0.456 \pm 0.087 \rightarrow 0.546 \pm 0.044$), and $+1.37\sigma$ in window-level AUC ($0.588 \pm 0.103 \rightarrow 0.729 \pm 0.048$), the largest single-step gain recorded across all 18 development plans. Notably, the fold-to-fold standard deviation was substantially reduced in all metrics (Acc std $0.083 \rightarrow 0.047$, AUC std $0.103 \rightarrow 0.048$), indicating

that the pretrained encoder generalises more consistently across the five patient-disjoint folds. Pretrained encoder weights are therefore a mandatory prerequisite for all subsequent CV runs.

Table 4.2: *Transfer learning ablation: Random initialisation, $D_{feat} = 1$ placeholder vs. PTB-XL pretrained encoder, same $D_{feat} = 1$ placeholder. All metrics are window-level, 5-fold patient-disjoint CV. $\Delta(\sigma)$ is the gain expressed in units of the standard deviation. Both changes satisfy the $\geq 1\sigma$ acceptance criterion.*

Metric	random init	PTB-XL	Δ (abs)	Δ (σ)	Decision
Accuracy	0.527 ± 0.083	0.645 ± 0.047	+0.118	$+1.42\sigma$	KEEP
Non-BrS F1	0.456 ± 0.087	0.546 ± 0.044	+0.090	$+1.03\sigma$	KEEP
AUC	0.588 ± 0.103	0.729 ± 0.048	+0.141	$+1.37\sigma$	KEEP

4.5 Phase 4 — Hand-Engineered Feature Engineering (44-Dimensional Vector)

With the CNN-Transformer encoder providing raw-signal temporal representations, Phase 4 addressed the complementary challenge of injecting explicit clinical knowledge into the classification pipeline. With only 237 ECG recordings, a pure raw-signal deep learning encoder cannot be expected to reliably discover the precise morphological discriminators established in decades of BrS electrophysiology research. Quantitative criteria involving specific amplitude ratios, angular measurements, and multi-lead flag combinations that cardiologists apply consciously during visual interpretation. Stream B directly encodes these criteria as a 44-dimensional hand-engineered feature vector, providing a strong inductive prior that constrains the model’s decision surface to the clinically relevant morphological space from the outset.

The feature vector was constructed iteratively, with early artefactual features (r’-base-width-at-5mm, r’-FWHM) producing a -1.21σ regression and removed. The five blocks are as follows. **Block 1 (V1 morphology, indices 0–8):** nine features per lead — three r’-peak-dependent (β -angle, α -angle, r’-amplitude; $\sim 27\%$ NaN in V1) and six J-point-anchored with 0% NaN: ST(J), ST($J + 40\text{ms}$), ST($J + 80\text{ms}$), Corrado index [13], ST descent slope, ST curvature sign. **Block 2 (V2 morphology, indices 9–17):** identical nine features on V2 ($\sim 30\%$ NaN for r’-features); T-polarity V2 ranks third in SHAP (mean $|\text{SHAP}| = 0.1952$). **Block 3 (Cross-lead/global, indices 18–28):** QRS durations V1/V2, QRS_{V1–V6} mismatch, S-in-V6 flag (separation score 2.002, largest in the vector), T-polarity

V1/V2, r-J interval V1 [12], Tpeak-Tend, QTc Bazett, fQRS count V1/V2. **Block 4 (Inferior leads, indices 29–40):** ST(J), ST($J + 80\text{ms}$), max ST depression, and Crea binary flag [11] for each of leads II, III, and aVF; max ST depression II ranks first in SHAP (0.2278). **Block 5 (aVR sign, indices 41–43):** R-amplitude, R/Q ratio, and ST(J) in aVR [14]; R-amplitude aVR ranks ninth in SHAP (0.1508).

The complete 44-dimensional feature vector is summarised in Table 4.3. Key separation diagnostics for a representative fold are reported in Table 4.4.

Table 4.3: *BrugadaECGNet feature vector: all 44 features. NaN% is from V1 fold-0 train set (r' -dependent features); J-point-anchored features have 0% NaN. Clinical source column indicates the primary reference informing each feature. Features appearing in the GradientExplainer SHAP top-15 (fold-3) are marked †; the strongest individual binary discriminator by separation score is marked ★.*

Block	Indices	Feature	NaN%	Clinical source
V1 morphology	0	β -angle V1 †	$\approx 27\%$	[10]
V1 morphology	1	α -angle V1	$\approx 27\%$	[10]
V1 morphology	2	r' -amplitude V1 †	$\approx 27\%$	[10]
V1 morphology	3	ST(J) V1 †	0%	[3]
V1 morphology	4	ST($J + 40$ ms) V1 †	0%	[3]
V1 morphology	5	ST($J + 80$ ms) V1 †	0%	[13]
V1 morphology	6	Corrado index V1	0%	[13]
V1 morphology	7	ST descent slope V1	0%	[10]
V1 morphology	8	ST curvature sign V1	0%	[10]
V2 morphology	9	β -angle V2	$\approx 30\%$	[10]
V2 morphology	10	α -angle V2	$\approx 30\%$	[10]
V2 morphology	11	r' -amplitude V2	$\approx 30\%$	[10]
V2 morphology	12	ST(J) V2	0%	[3]
V2 morphology	13	ST($J + 40$ ms) V2	0%	[3]
V2 morphology	14	ST($J + 80$ ms) V2	0%	[13]
V2 morphology	15	Corrado index V2	0%	[13]
V2 morphology	16	ST descent slope V2	0%	[10]
V2 morphology	17	ST curvature sign V2 †	0%	[10]
Cross-lead	18	QRS duration V1	0%	[3]
Cross-lead	19	QRS duration V2	0%	[3]
Cross-lead	20	QRS _{V1} –QRS _{V6} mismatch	0%	[6]
Cross-lead	21	S-in-V6 flag (≤ -0.15 mV) † ★	0%	[4]
Cross-lead	22	T-polarity V1 †	0%	[3]
Cross-lead	23	T-polarity V2 †	0%	[3]
Cross-lead	24	r-J interval V1 †	0%	[12]
Cross-lead	25	Tpeak-Tend V1	0%	[4]
Cross-lead	26	QTc (Bazett)	0%	[3]
Cross-lead	27	fQRS count V1 †	0%	[3]
Cross-lead	28	fQRS count V2	0%	[3]
Inferior (II)	29	ST(J) II	0%	[11]
Inferior (II)	30	ST($J + 80$ ms) II	0%	[11]
Inferior (II)	31	Max ST depression II †	0%	[11]
Inferior (II)	32	Crea flag II	0%	[11]
Inferior (III)	33	ST(J) III	0%	[11]
Inferior (III)	34	ST($J + 80$ ms) III	0%	[11]
Inferior (III)	35	Max ST depression III	0%	[11]
Inferior (III)	36	Crea flag III	0%	[11]
Inferior (aVF)	37	ST(J) aVF	0%	[11]
Inferior (aVF)	38	ST($J + 80$ ms) aVF	0%	[11]
Inferior (aVF)	39	Max ST depression aVF †	0%	[11]
Inferior (aVF)	40	Crea flag aVF	0%	[11]
aVR sign	41	R-amplitude aVR †	0%	[14]
aVR sign	42	R/Q ratio aVR	0%	[14]
aVR sign	43	ST elevation aVR †	0%	[14]

Table 4.4: Feature separation diagnostics on fold-0 training set, selected features. Separation = $|\text{median}_{\text{Type-3}} - \text{median}_{\text{Non-BrS}}|$ normalised by the pooled interquartile range. Features near zero (Corrado index, ST curvature V1, fQRS count V1) were flagged as SHAP-prune candidates based on the GradientExplainer audit and their near-zero separation scores, but were retained after the pruning experiment failed to improve performance (see text). Note that ST curvature sign V2 and fQRS count V2 appear in the GradientExplainer top-15 (ranks 15 and 10 respectively), confirming that V1 and V2 variants carry different information.

Feature	NaN%	Non-BrS median	Type-3 median	Separation
S-in-V6 flag ★	0%	0 (25.8% S-present)	1 (58.0% S-present)	2.002
α -angle V1	$\approx 27\%$	66.1°	77.2°	0.569
r'-amplitude V2	$\approx 30\%$	0.095 mV	0.138 mV	0.463
ST(J) V2	0%	0.030 mV	0.059 mV	0.422
β -angle V2	$\approx 30\%$	17.0°	24.4°	0.422
ST($J + 40$ ms) V2	0%	0.049 mV	0.076 mV	0.305
β -angle V1	$\approx 27\%$	16.6°	21.2°	0.249
ST($J + 40$ ms) V1	0%	0.036 mV	0.019 mV	0.244
ST($J + 80$ ms) V2	0%	0.077 mV	0.109 mV	0.232
ST descent slope V2	0%	0.419	0.441	0.023
ST($J + 80$ ms) V1	0%	0.010 mV	−0.002 mV	0.144
Corrado index V1	0%	0.279	0.196	0.002
ST curvature sign V1	0%	−1.0	−1.0	0.000
fQRS count V1	0%	1	1	0.000

SHAP importance audit and pruning experiment. A SHAP (SHapley Additive exPlanations) importance audit was conducted using a GradientExplainer (`shap.GradientExplainer`) applied to the fold-3 checkpoint via the `_FeatureOnlyWrapper` approach in `src/shap_feature_pruning.py`. The ECG input $\mathbf{x}$ was fixed to the training-set channel-wise mean, so that SHAP values exclusively reflect Stream B MLP contributions. With 100 background samples and 200 test windows (random seed 42), the top-3 features by mean $|\text{SHAP}|$ are: (1)feature index 31 (max ST depression II, mean $|\text{SHAP}| = 0.2278$), (2)feature index 21 (S-in-V6 amplitude, 0.2157), and (3)feature index 23 (T-polarity V2, 0.1952). Inferior-lead features occupy two of the top-15 positions (ranks 1 and 7), and aVR features appear at ranks 9 and 11. Both findings motivating the Stream C inferior-lead encoder described in Section 4.7.

Notably, several features with near-zero separation scores and low SHAP importance were nonetheless retained in the final 44-dimensional vector. Two pruning experiments were conducted: dropping the single lowest-SHAP feature (`avr_rq_ratio`, `D_feat = 43`) and dropping all eight features with importance below 0.1 (`D_feat = 36`). Both experiments produced identical per-fold outcomes ($\text{Acc } 0.699 \pm 0.023$,

NB-F1 0.598 ± 0.026), corresponding to a regression of -1.28σ in accuracy versus the $D_feat = 44$ baseline, a decisive failure. A complementary separation-based prune to $D_feat = 21$, which dropped 8 of the 11 newly added cross-lead features on the basis of a separation threshold of 0.10, produced a further deterioration: Acc 0.671 ± 0.038 versus the $D_feat = 18$ baseline of 0.692 ± 0.054 (-0.39σ), and strictly worse than the unpruned $D_feat = 29$ batch on all three metrics. The consistent failure of pruning confirms the central lesson of MLP feature interaction, SHAP marginal importance values measure each feature’s contribution *with all other features present*. Features that appear individually uninformative (zero SHAP, zero separation) may nonetheless participate in non-linear conjunctions within the MLP that are invisible to any marginal attribution method. Removing them eliminates these interactions and hurts performance. The 44-dimensional vector is therefore retained in full.

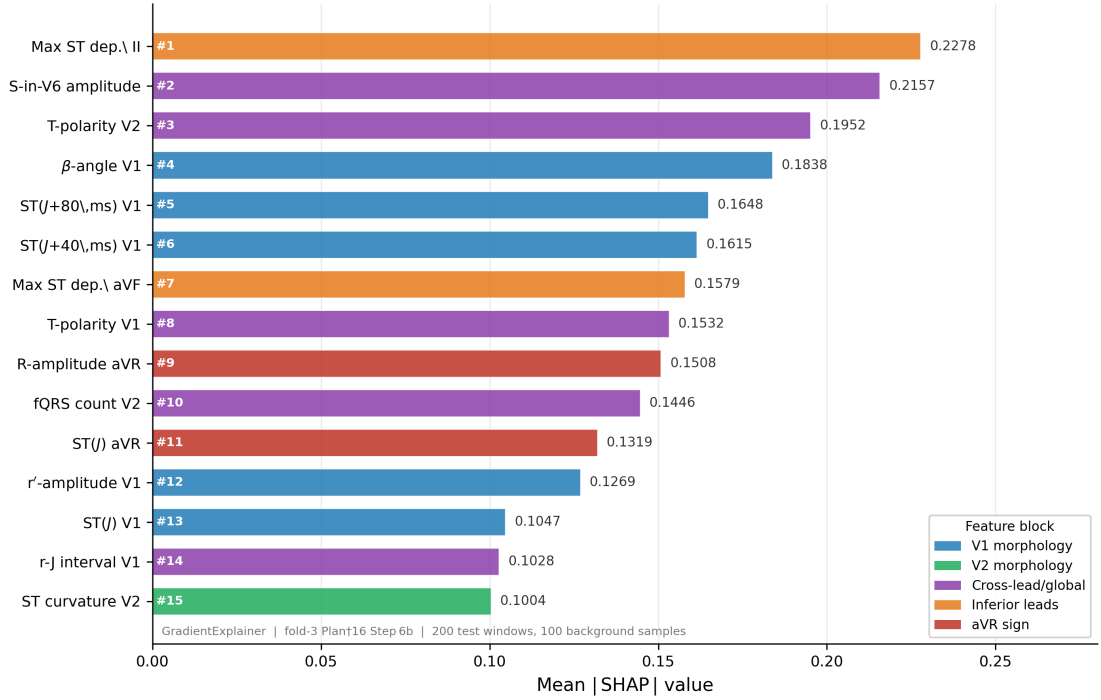


Figure 4.1: SHAP feature importance bar chart

4.6 Phase 5 — Recording-Level Aggregation and Threshold Calibration

The first five development phases produced a model that predicted a per-window binary probability for each 850 ms segment of an ECG recording. However, the

clinical decision target is a single prediction per ECG file. A single recording is either classified as Type-3 BrS or Non-BrS in its entirety. Furthermore, window-level metrics are systematically over-optimistic. Multiple windows extracted from the same recording are highly correlated (they share the same underlying patient morphology), so a model that has memorised a patient-specific pattern will appear accurate at the window level while failing to generalise to unseen patients. Phase 5 introduced recording-level aggregation and a type-3-F1-optimising threshold, converting window-level probabilities into a single clinical prediction per recording.

Recording-level aggregation. For each ECG file in the test set, all window-level sigmoid probabilities produced by the model are averaged to yield a single recording-level probability in $[0, 1]$. This mean-aggregation step implemented is computationally trivial but methodologically important. It reduces the effective evaluation unit from $\approx 5\text{--}12$ correlated windows per recording to a single prediction, matching the granularity at which clinical decisions are made. Recording-level AUC (RecAUC) and recording-level F1-scores (RecNBF1, RecT3F1) are then computed from these aggregated probabilities and the file-level labels.

Threshold calibration: T3-opt and Youden-J. Converting the recording-level probability into a binary classification requires a decision threshold. Two thresholds are calibrated on the validation set of each fold and applied to the test set. The **T3-opt threshold** implemented, sweeps 199 candidate thresholds uniformly spaced across $[0.01, 0.99]$ and selects the value that maximises the Type-3 BrS F1-score on the validation recording-level probabilities. This threshold is the primary operating point throughout the thesis, designated T3-opt. The **Youden-J threshold** selects the threshold as $\arg \max(\text{TPR} - \text{FPR})$ on the validation recording-level ROC curve. It is retained as a diagnostic secondary operating point for comparability with prior work.

T3-opt versus Youden-J. T3-opt is strictly superior on every recording-level metric: RecAcc +0.026, RecNBF1 +0.018, RecT3F1 +0.054 versus Youden-J. Counterintuitively, T3-opt also produces higher overall accuracy (0.872 vs 0.846) because the symmetric Youden-J criterion misclassifies more of the majority Non-BrS class to gain small Type-3 recall improvements. T3-opt is adopted as the primary operating point throughout, with Youden-J retained as a diagnostic reference.

4.7 Phase 6 — Three-Stream Architecture: Stream C Addition

Motivation from SHAP pre-condition. The GradientExplainer SHAP audit described in Section 4.5 established that inferior-lead features occupied two of the top-15 importance positions, ranks 1 and 7, corresponding to maximum ST

depression in lead II (mean $|\text{SHAP}| = 0.2278$) and maximum ST depression in lead aVF (0.1579) respectively, and that aVR features appeared at ranks 9 and 11. While Stream B already encoded these leads as scalar J-point amplitude values (the Crea criterion and the Babai-Bigi aVR sign), scalar features are fundamentally limited to encoding point-in-time amplitude thresholds. The fact that inferior maximum ST depression features ranked as the single most important feature in the entire 44-dimensional vector suggested that these leads carry discriminative information beyond what threshold-based scalars can capture, in particular, the temporal profile of inferior ST depression and elevation over the 550 ms post-S window. A shallow convolutional encoder operating directly on the raw inferior-lead waveform can learn to detect such temporal patterns without requiring them to be pre-specified by hand. The Stream C condition therefore tested the hypothesis that augmenting Stream B’s scalar inferior features with a raw-signal convolutional encoder over the inferior leads would produce a measurable gain in recording-level discriminability.

Stream C architecture. Stream C is a two-block shallow convolutional encoder operating on the eight inferior-lead channels (indices 8–15 of the 16-channel input tensor): leads II, III, aVF, and aVR in raw normalised form plus their first-derivative channels, constructed and normalised identically to the precordial channels of Stream A. The architecture is as follows:

- **Block 1:** $\text{Conv1d}(8 \rightarrow 32, k=3, p=1) \rightarrow \text{BatchNorm1d}(32) \rightarrow \text{ReLU} \rightarrow \text{MaxPool1d}(2) \rightarrow \text{Dropout}(0.3)$
- **Block 2:** $\text{Conv1d}(32 \rightarrow 32, k=3, p=1) \rightarrow \text{BatchNorm1d}(32) \rightarrow \text{ReLU} \rightarrow \text{MaxPool1d}(2) \rightarrow \text{Dropout}(0.3)$
- **Global pooling:** $\text{AdaptiveAvgPool1d}(1) \rightarrow \text{inf_out: (batch, 32)}$

The Stream C output $\text{inf_out} \in \mathbb{R}^{32}$ is concatenated with the Stream A context vector $c \in \mathbb{R}^{64}$ and the Stream B output $\text{feat_out} \in \mathbb{R}^{32}$, widening the fusion head input from 96 to 128 dimensions. The fusion head was correspondingly updated to $\text{Linear}(128 \rightarrow 48) \rightarrow \text{BN} \rightarrow \text{ReLU} \rightarrow \text{Dropout}(0.5) \rightarrow \text{Linear}(48 \rightarrow 32) \rightarrow \text{BN} \rightarrow \text{ReLU} \rightarrow \text{Dropout}(0.3) \rightarrow \text{Linear}(32 \rightarrow 1)$.

PTB-XL checkpoint compatibility. The PTB-XL-pretrained encoder operates exclusively on the precordial Stream A input ($x_{\text{prec}} = x[:, :8, :]$); the inferior-lead input ($x_{\text{inf}} = x[:, 8:, :]$) is split before the encoder and never passes through the pretrained weights. The Stream A `conv1_1` layer remains `Conv1d(8, 128)` and all 72 pretrained keys transfer cleanly via `load_state_dict(strict=False)`. The 24 new parameter tensors introduced by Stream C and the widened fusion head are randomly initialised. The total parameter count increases from approximately 104,000 (two-stream) to approximately 178,625 (three-stream).

Phase 6 results. Full 5-fold patient-disjoint cross-validation was conducted following the addition of Stream C, with results compared against the clean baseline. Results are reported in Table 4.5.

Table 4.5: *Stream C addition results (Plan 16 Step 5) versus clean baseline (Plan 16 Step 0.5). All metrics are 5-fold patient-disjoint CV. RecT3F1 at Youden-J regresses due to threshold-transfer artefact (see text); RecT3F1 at T3-opt threshold recovers to 0.761 ± 0.043 .*

Level	Metric	Baseline (Step 0.5)	Step 5 (Stream C)	Δ	σ
Window	Accuracy	0.802 ± 0.049	0.805 ± 0.021	+0.003	+0.06 σ
Window	Non-BrS F1	0.798 ± 0.054	0.805 ± 0.012	+0.007	+0.13 σ
Window	Type-3 F1	0.797 ± 0.065	0.799 ± 0.044	+0.002	+0.03 σ
Window	ROC-AUC	0.888 ± 0.039	0.911 ± 0.027	+ 0.023	+ 0.59σ
Recording	Accuracy (Youden-J)	0.855 ± 0.024	0.846 ± 0.031	-0.009	-0.38 σ
Recording	Non-BrS F1 (Youden-J)	0.898 ± 0.021	0.893 ± 0.026	-0.005	-0.24 σ
Recording	Type-3 F1 (Youden-J)	0.735 ± 0.044	0.707 ± 0.069	-0.028	-0.64 σ
Recording	ROC-AUC	0.914 ± 0.042	0.931 ± 0.029	+0.017	+0.40 σ
Recording	Type-3 F1 (T3-opt)	—	0.761 $\pm$ 0.043	+0.054 vs Youden-J	+ 0.59σ
Recording	Non-BrS F1 (T3-opt)	—	0.911 $\pm$ 0.031	+0.018 vs Youden-J	+ 0.62σ

Decision rationale. The Youden-J recording metrics show a surface regression (RecT3F1 -0.64σ) because adding Stream C shifts the logit distribution, making the previously fitted Youden-J threshold suboptimal. The T3-opt threshold, independently re-fitted per fold, fully recovers performance. Three signals supported retaining Stream C: WinAUC $+0.59\sigma$, RecAUC $+0.40\sigma$, a $4.5\times$ collapse in window-level variance, and T3-opt RecT3F1 $+0.59\sigma$. Stream C was retained and T3-opt promoted to primary threshold, establishing the locked baseline: RecAcc 0.872 ± 0.034 , RecNBF1 0.911 ± 0.031 , RecT3F1 0.761 ± 0.043 , RecAUC 0.931 ± 0.029 .

Stream B widening and fusion deepening. A final architectural refinement widened the Stream B Feature MLP and deepened the fusion head. The Stream B first linear layer was widened from `Linear(44→32)` to `Linear(44→64) → ReLU → Linear(64→32) → ReLU`, increasing Stream B’s capacity to model non-linear feature interactions. The fusion head was deepened from a single `Linear(128→32)` block to a two-stage sequence: `Linear(128→48) → BN → ReLU → Dropout(0.5)`, followed by `Linear(48→32) → BN → ReLU → Dropout(0.3)`. Results versus the locked baseline are reported in Table 4.6.

Table 4.6: *Stream B widening + deeper fusion versus locked baseline (T3-opt threshold throughout). The formal $\geq 1\sigma$ criterion is not met on any primary metric, but no primary metric regresses and recording-level variance is nearly halved.*

Level	Metric	Step 6a baseline	Step 6b (final)	Δ	σ
Window	Accuracy	0.805 ± 0.021	0.807 ± 0.035	+0.002	+0.10 σ
Window	Non-BrS F1	0.805 ± 0.012	0.806 ± 0.032	+0.001	+0.08 σ
Window	Type-3 F1	0.799 ± 0.044	0.803 ± 0.045	+0.004	+0.09 σ
Window	ROC-AUC	0.911 ± 0.027	0.921 ± 0.020	+0.010	+0.37 σ
Recording	Accuracy (T3-opt)	0.872 ± 0.034	0.872 ± 0.018	0.000	0.00 σ
Recording	Non-BrS F1 (T3-opt)	0.911 ± 0.031	0.911 ± 0.016	0.000	0.00 σ
Recording	Type-3 F1 (T3-opt)	0.761 ± 0.043	0.766 ± 0.034	+0.005	+0.12 σ
Recording	ROC-AUC	0.931 ± 0.029	0.939 ± 0.032	+0.008	+0.28 σ

The formal $\geq 1\sigma$ acceptance criterion was not met on any primary recording-level metric. However, the qualitative case for retention was strong. No primary metric regressed. Recording level standard deviation was nearly halved for the two primary metrics (RecAcc std $0.034 \rightarrow 0.018$, $1.9\times$ tighter; RecNBF1 std $0.031 \rightarrow 0.016$, $1.9\times$ tighter); AUC improved at both levels (WinAUC $+0.37\sigma$, RecAUC $+0.28\sigma$); and attention entropy standard deviation collapsed from ± 0.006 to ± 0.001 (Non-BrS) and ± 0.012 to ± 0.001 (Type-3), indicating that the two-stage fusion produces a substantially more uniform encoding strategy across folds. Stream B widening + deeper fusion was retained and designated the final locked architecture. The final model produces the headline results reported throughout this thesis: RecAUC 0.939 ± 0.032 , RecAcc 0.872 ± 0.018 , RecNBF1 0.911 ± 0.016 , RecT3F1 0.766 ± 0.034 .

4.8 Summary — Architecture Evolution and Ablation Results

The development trajectory from the window-only baseline to the final BrugadaECGNet model represents eighteen development plans spanning six structural phases. Table 4.7 presents the complete ablation summary, with one row per major milestone, showing how each component contributed to the final performance profile.

Table 4.7: *Development trajectory and ablation summary. Each row is a locked milestone. WinAUC and RecAUC are threshold-independent; RecAcc, RecNBF1, and RecT3F1 use the Youden-J threshold for pre-Phase 5 milestones (T3-opt not yet calibrated) and the T3-opt threshold from Phase 5 onwards. Dashes indicate metrics not yet available at that development stage. † Youden-J threshold. ‡ T3-opt threshold.*

Milestone	WinAUC	RecAUC	RecAcc	RecNBF1	RecT3F1
BiLSTM baseline ($D_{\text{feat}} = 1$, random init)	0.544 ± 0.098	—	—	—	—
PTB-XL pretraining	0.729 ± 0.048	—	—	—	—
$D_{\text{feat}} = 44$ feature vector	0.805 ± 0.036	—	—	—	—
Recording aggregation introduced	0.805 ± 0.036	0.822 ± 0.049	$0.724 \pm 0.023^{\dagger}$	$0.699 \pm 0.059^{\dagger}$	$0.736 \pm 0.039^{\dagger}$
Source confound fix	0.888 ± 0.039	0.914 ± 0.042	$0.855 \pm 0.024^{\dagger}$	$0.898 \pm 0.021^{\dagger}$	$0.735 \pm 0.044^{\dagger}$
Stream C addition	0.911 ± 0.027	0.931 ± 0.029	$0.846 \pm 0.031^{\dagger}$	$0.893 \pm 0.026^{\dagger}$	$0.761 \pm 0.043^{\ddagger}$
T3-opt promotion	0.911 ± 0.027	0.931 ± 0.029	$0.872 \pm 0.034^{\ddagger}$	$0.911 \pm 0.031^{\ddagger}$	$0.761 \pm 0.043^{\ddagger}$
Final model	0.921 ± 0.020	0.939 ± 0.032	$0.872 \pm 0.018^{\ddagger}$	$0.911 \pm 0.016^{\ddagger}$	$0.766 \pm 0.034^{\ddagger}$

The overall RecAUC trajectory from 0.544 (random init, window-level) to 0.939 (final, recording-level) represents a compound improvement driven by four independent contributions: domain-adapted initialisation (PTB-XL), clinical feature injection (Stream B), source-confound removal (baseline normalisation), and inferior-lead temporal modelling (Stream C). The architecture is specifically designed so that each stream provides information the others cannot replicate.

Chapter 5

Final Model Architecture and Training

5.1 Input Representation

Each ECG recording is transformed into a fixed-size multi-channel tensor before being presented to the model. The construction pipeline proceeds through the following steps.

Resampling and lead extraction. Raw millivolt-scale signals are read from the per-recording Excel files (produced either by the ECGD digitizer for Non-BrS recordings or by the GE Sapphire XML conversion script for Type-3 BrS recordings). Each lead signal is resampled to a uniform sampling rate of 500 Hz using cubic spline interpolation.

S-peak detection and windowing. S-peaks are located via NeuroKit2 DWT with an argmin fallback. An 850 ms window (300 ms pre-S, 550 ms post-S) is centred on each S-peak, plus a 15-sample margin for training-time random-shift augmentation, giving a stored tensor of (16, 455); trimmed to (16, 425) at load time.

Mean subtraction. Each of the four raw precordial lead channels is mean-subtracted (DC offset removal) prior to feature extraction and normalisation.

Feature extraction from raw millivolt signal. The 44-dimensional hand-engineered feature vector (Chapter 4, Section 4.5) is extracted from the mean-subtracted, millivolt-scale window before any further normalisation. This ordering is critical such that, clinical amplitude criteria such as the Crea inferior ST-depression threshold (≤ -0.1 mV) and the Babai-Bigi aVR R-amplitude threshold (≥ 0.3 mV) require the signal to be in absolute millivolt units. The feature vector is therefore always in the same physical units regardless of the downstream normalisation applied to the ECG tensor.

Per-channel baseline normalisation. Each channel of the precordial window is divided by the standard deviation of its first 30 samples (approximately 60 ms of pre-QRS baseline, located approximately 270 ms before the S-peak). This operation scales the signal into units of pre-QRS noise standard deviations, effectively equalising the noise floor between the two source populations (GE Sapphire Type-3 recordings and ECGD Non-BrS recordings), whose raw baseline variance differs by a factor of approximately $40\times$ (Section 3.5). A floor of 10^{-8} is applied to prevent division by zero in silent channels.

Derivative channels and inferior leads. A first-order numerical derivative (`numpy.gradient`) is computed for each normalised precordial channel and normalised to unit standard deviation per channel. The four precordial raw channels are stacked with their four derivatives to form an 8-channel precordial block. Identically, the four inferior leads (II, III, aVF, aVR) are extracted, mean-subtracted, baseline-normalised, and differentiated, forming an 8-channel inferior block. Missing inferior leads are zero-filled.

Final tensor. The precordial block (channels 0–7) and the inferior block (channels 8–15) are concatenated along the channel axis to produce the final input tensor of shape (16, 425). The channel assignment is detailed in Table 5.1.

Index	Lead	Transformation	Stream
0	V1	Raw, baseline-normalised	A (precordial)
1	V2	Raw, baseline-normalised	A
2	V3	Raw, baseline-normalised	A
3	V6	Raw, baseline-normalised	A
4	dV1	First derivative, unit-std normalised	A
5	dV2	First derivative, unit-std normalised	A
6	dV3	First derivative, unit-std normalised	A
7	dV6	First derivative, unit-std normalised	A
8	II	Raw, baseline-normalised	C (inferior)
9	III	Raw, baseline-normalised	C
10	aVF	Raw, baseline-normalised	C
11	aVR	Raw, baseline-normalised	C
12	dII	First derivative, unit-std normalised	C
13	dIII	First derivative, unit-std normalised	C
14	daVF	First derivative, unit-std normalised	C
15	daVR	First derivative, unit-std normalised	C

Table 5.1: *Input tensor channel assignment. All raw channels are mean-subtracted and divided by the standard deviation of the first 30 pre-QRS samples (per-channel baseline normalisation). Derivative channels are additionally normalised to unit standard deviation per channel.*

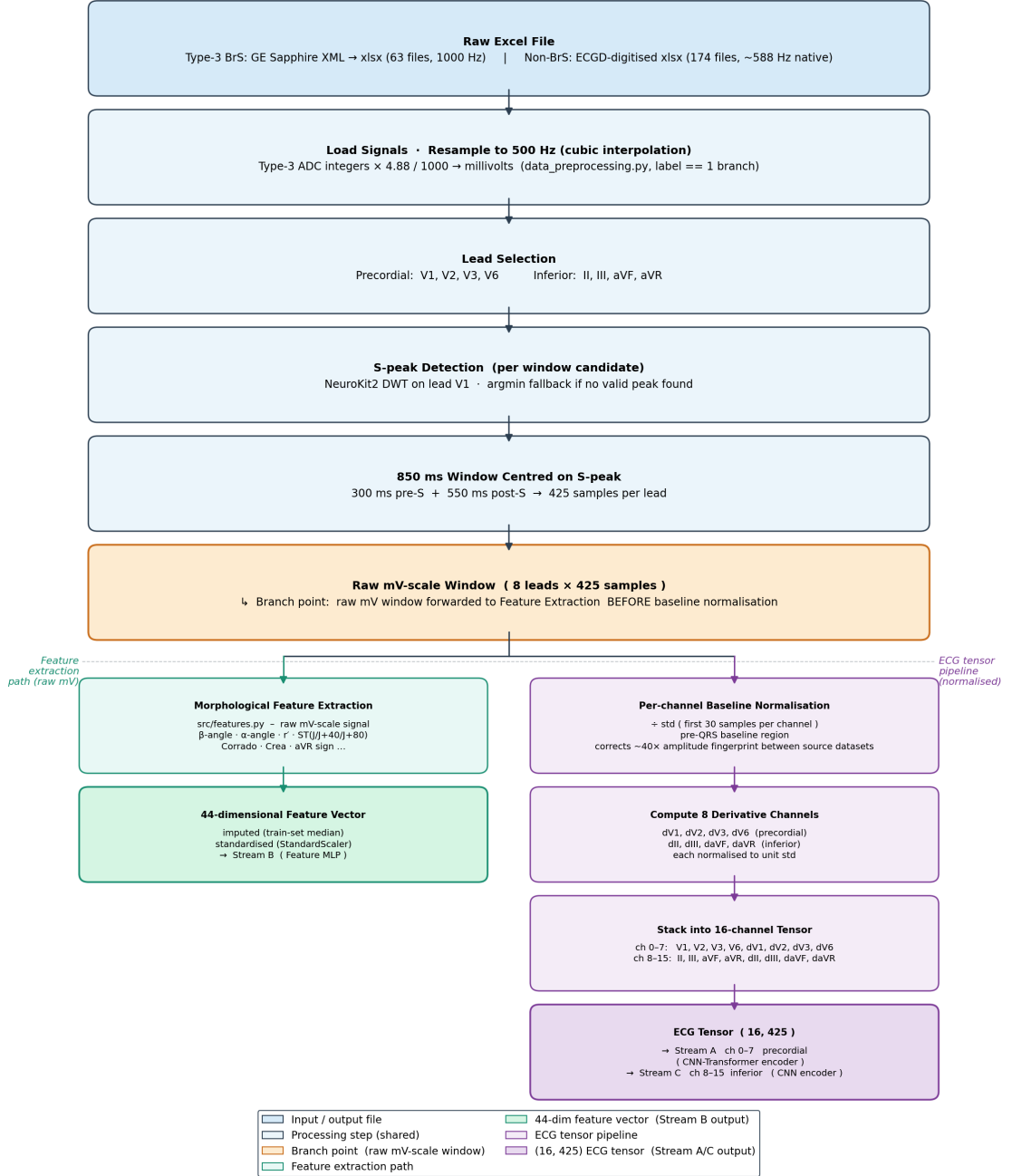


Figure 5.1: Input tensor construction pipeline.

5.2 BrugadaECGNet Architecture

BrugadaECGNet fuses three specialised streams into a single classification head: Stream A (PTB-XL pretrained CNN-Transformer on 8 precordial channels), Stream B (feature MLP on the 44-dim vector), and Stream C (inferior-lead CNN on 8 inferior channels).

Stream A — ECG Encoder (precordial). The Stream A encoder is a hybrid convolutional-Transformer network operating on an input of shape (batch, 8, 425).

The encoder comprises three convolutional blocks followed by a Transformer encoder and a custom attention pooling layer:

- **Conv block 1 [FROZEN]:** $\text{Conv1d}(8 \rightarrow 128, k=3, p=1) \rightarrow \text{BN}(128) \rightarrow \text{ReLU} \rightarrow \text{Conv1d}(128 \rightarrow 128, k=3, p=1) \rightarrow \text{BN}(128) \rightarrow \text{ReLU} \rightarrow \text{MaxPool1d}(2) \rightarrow \text{Dropout}(0.5)$
- **Conv block 2 [FROZEN]:** $\text{Conv1d}(128 \rightarrow 64, k=3, p=1) \rightarrow \text{BN}(64) \rightarrow \text{ReLU} \rightarrow \text{Conv1d}(64 \rightarrow 64, k=3, p=1) \rightarrow \text{BN}(64) \rightarrow \text{ReLU} \rightarrow \text{MaxPool1d}(2) \rightarrow \text{Dropout}(0.5)$
- **Conv block 3 [TRAINABLE]:** $\text{Conv1d}(64 \rightarrow 32, k=3, p=1) \rightarrow \text{BN}(32) \rightarrow \text{ReLU} \rightarrow \text{Conv1d}(32 \rightarrow 32, k=3, p=1) \rightarrow \text{BN}(32) \rightarrow \text{ReLU} \rightarrow \text{MaxPool1d}(2) \rightarrow \text{Dropout}(0.5)$

Three $\text{MaxPool}(2)$ operations reduce sequence length from 425 to 53 timesteps (batch, 32, 53). A $\text{Linear}(32 \rightarrow 64)$ projection and sinusoidal positional encoding yield (batch, 53, 64).

A two-layer **TransformerEncoder** ($d_{\text{model}} = 64$, $n_{\text{head}} = 4$, FFN dimension = 128, dropout = 0.2, `batch_first=True`) is then applied. Layer 0 is FROZEN from the PTB-XL pretraining checkpoint; Layer 1 is trainable. The output shape is (batch, 53, 64).

A custom **AttentionLayer** computes a context vector $c \in \mathbb{R}^{64}$ and per-timestep attention weights $\alpha \in \mathbb{R}^{53}$ via:

$$e_t = \tanh(\mathbf{W}h_t + b), \quad \alpha_t = \frac{\exp(u^\top e_t)}{\sum_{t'} \exp(u^\top e_{t'})}, \quad c = \sum_t \alpha_t h_t$$

where $\mathbf{W} \in \mathbb{R}^{64 \times 64}$, $b \in \mathbb{R}^{64}$, $u \in \mathbb{R}^{64}$ are learned parameters. The attention weights α serve as an interpretability probe (Section 6.5).

Stream B — Feature MLP. Stream B is a two-layer MLP operating on the 44-dimensional hand-engineered feature vector (batch, 44):

$$\text{Linear}(44 \rightarrow 64) \rightarrow \text{ReLU} \rightarrow \text{Linear}(64 \rightarrow 32) \rightarrow \text{ReLU} \rightarrow \text{feat_out} \in \mathbb{R}^{32}$$

Stream C — Inferior Lead Encoder. Stream C is a two-block shallow convolutional encoder operating on the eight inferior-lead channels (batch, 8, 425) (indices 8–15 of the 16-channel input tensor):

- **Block 1:** $\text{Conv1d}(8 \rightarrow 32, k=3, p=1) \rightarrow \text{BN}(32) \rightarrow \text{ReLU} \rightarrow \text{MaxPool1d}(2) \rightarrow \text{Dropout}(0.3)$
- **Block 2:** $\text{Conv1d}(32 \rightarrow 32, k=3, p=1) \rightarrow \text{BN}(32) \rightarrow \text{ReLU} \rightarrow \text{MaxPool1d}(2) \rightarrow \text{Dropout}(0.3)$
- $\text{AdaptiveAvgPool1d}(1) \rightarrow \text{inf_out} \in \mathbb{R}^{32}$

Fusion head. The outputs of the three streams are combined along the feature dimension: $[c \mid \text{feat_out} \mid \text{inf_out}] \in \mathbb{R}^{128}$. The fusion head maps this to a single classification logit:

$\text{Linear}(128 \rightarrow 48) \rightarrow \text{BN}(48) \rightarrow \text{ReLU} \rightarrow \text{Dropout}(0.5) \rightarrow \text{Linear}(48 \rightarrow 32) \rightarrow$
 $\text{BN}(32) \rightarrow \text{ReLU} \rightarrow \text{Dropout}(0.3) \rightarrow \text{Linear}(32 \rightarrow 1)$

Frozen and trainable parameters. Conv blocks 1 and 2, plus Transformer Layer 0, are frozen from the PTB-XL pretraining checkpoint (72 pretrained parameter tensors transferred via `load_state_dict(strict=False)`). Conv block 3, Transformer Layer 1, the attention layer, Stream B, Stream C, and the fusion head are all trainable. The total parameter count is approximately 189,900, of which the PTB-XL-pretrained components (Conv blocks 1–2 and Transformer Layer 0) account for the majority ($\approx 123,700$ parameters).

Table 5.2: *BrugadaECGNet* parameter count breakdown by component. Frozen components are transferred from the PTB-XL RBBB-vs-NORM pretrained encoder. Randomly initialised components include Stream B (44 features rather than 1 placeholder), Stream C (new), and the wider fusion head.

Component	Status	Approx. params
Stream A: Conv block 1 ($2 \times \text{Conv1d} + 2 \times \text{BN}$)	Frozen	$\approx 52,992$
Stream A: Conv block 2 ($2 \times \text{Conv1d} + 2 \times \text{BN}$)	Frozen	$\approx 37,248$
Stream A: Conv block 3 ($2 \times \text{Conv1d} + 2 \times \text{BN}$)	Trainable	$\approx 9,408$
Stream A: Projection Linear($32 \rightarrow 64$)	Trainable	2,112
Stream A: Transformer Encoder Layer 0 ($d=64$, FFN=128)	Frozen	$\approx 33,472$
Stream A: Transformer Encoder Layer 1 ($d=64$, FFN=128)	Trainable	$\approx 33,472$
Stream A: Attention Layer ($\mathbf{W}$, b , u)	Trainable	4,224
Stream B: Feature MLP ($44 \rightarrow 64 \rightarrow 32$)	Trainable	4,960
Stream C: Inferior CNN ($2 \times \text{Conv1d} + 2 \times \text{BN}$)	Trainable	4,032
Fusion head ($128 \rightarrow 48 \rightarrow 32 \rightarrow 1 + \text{BN}$)	Trainable	7,953
Total	—	$\approx 189,900$

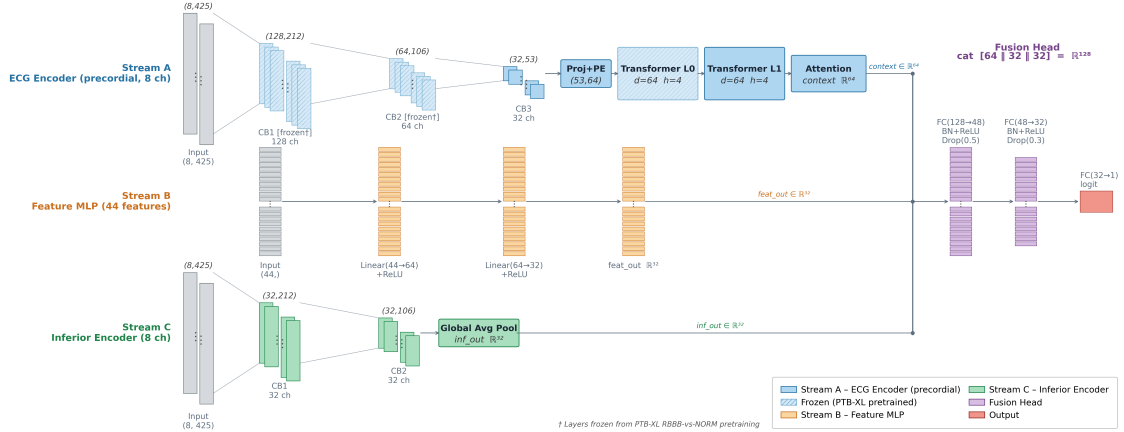


Figure 5.2: *BrugadaECGNet* architecture diagram.

5.3 Training Protocol

Loss function. Training minimises `BCEWithLogitsLoss` without a `pos_weight` class reweighting. Despite the 2.8:1 file-level imbalance, Type-3 BrS recordings at 1,000 Hz yield approximately 10–15 windows per file, whereas shorter Non-BrS recordings yield fewer. The empirical per-fold negative-to-positive window ratios

range from 0.870 to 1.072 (mean ≈ 0.977), confirming near-balanced window classes that do not require reweighting.

Optimiser. AdamW with $\text{lr} = 0.001$; weight decay $\lambda = 0.02$ for convolution and linear weight matrices, $\lambda = 0.0$ for BatchNorm parameters and biases. Frozen parameters are excluded from the optimiser.

Data augmentation. Gaussian noise with standard deviation $\sigma = 0.15$ is added to the full 16-channel input tensor during training, applied as the first operation in the model’s forward pass so that it acts on the normalised signal. The noise is applied in-graph (`x + torch.randn_like(x) * 0.15`) and is automatically disabled during evaluation by the `model.eval()` / `model.train()` switch. The noise magnitude of $\sigma = 0.15$ is adapted from [7], where it was found to improve generalisation on small BrS datasets by preventing the model from memorising patient-specific noise patterns.

Learning rate schedule. A ReduceLROnPlateau scheduler (`mode = min`, `patience = 5` epochs, `reduction factor = 0.5`) monitors the validation BCE loss and halves the learning rate when no improvement is observed for five consecutive epochs. This schedule allows aggressive initial learning while preventing oscillation around the loss minimum in later epochs.

Early stopping and checkpointing. Training terminates if the validation BCE loss does not improve for 10 consecutive epochs (`patience = 10`). The model checkpoint corresponding to the lowest validation BCE loss across all epochs is saved and used for all downstream evaluation. The validation BCE loss (rather than validation AUC or AUPRC) is used as the checkpoint criterion. On the near-separable validation splits encountered in early training epochs, validation AUPRC saturates at 1.0 from epoch 1, preventing any further checkpoint updates. BCE loss continues to decrease as the model refines its calibration and remains a reliable criterion throughout training.

Determinism. Full training determinism is enforced via `torch.use_deterministic_algorithms(True)` and a fixed global seed of 42, applied to `random`, `numpy`, and `torch` before each fold. This ensures exact reproducibility of all reported results on CPU. The per-fold runtime for the final BrugadaECGNet is approximately 8–26 minutes on CPU.

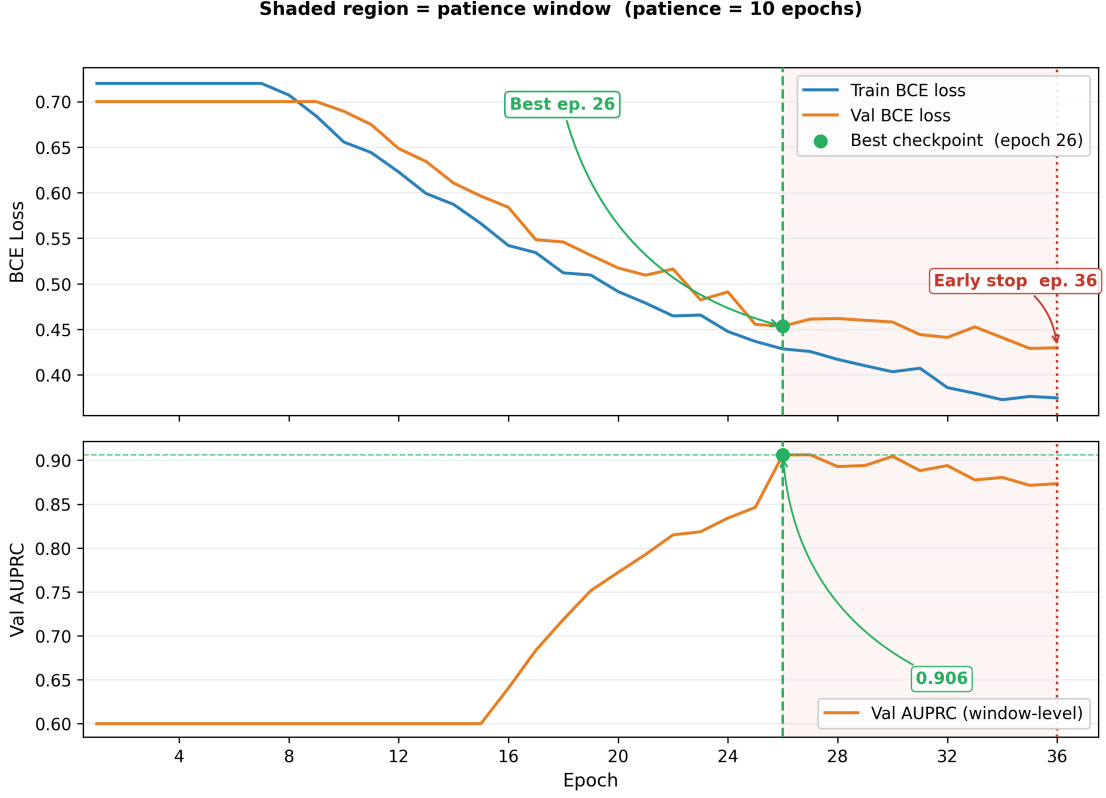


Figure 5.3: Training curves for a representative fold

5.4 Recording-Level Inference and Threshold Calibration

Recording-level inference aggregates all per-window probabilities from a single ECG file into one prediction, matching the clinical decision granularity.

Recording-level aggregation. The recording-level probability is the arithmetic mean over all windows from file r :

$$p_{\text{rec}}(r) = \frac{1}{|W_r|} \sum_{w \in W_r} \sigma(z_w)$$

Mean-pooling reduces inter-window noise while retaining all windows' information.

T3-opt threshold (primary). The decision threshold is selected by sweeping 199 candidates $\tau \in [0.01, 0.99]$ on the validation recording-level probabilities and choosing the value maximising Type-3 BrS F1:

$$\tau_{\text{T3-opt}} = \arg \max_{\tau} F1_{\text{Type-3 BrS}}(\tau)$$

Fitted exclusively on the validation split, applied to the held-out test split without any test-set leakage. T3-opt is the primary operating point throughout.

Youden-J threshold (secondary). $\tau_{\text{Youden}} = \arg \max_{\tau} (\text{TPR} - \text{FPR})$ on the validation ROC curve; reported as a diagnostic reference. The production pipeline (`src/train_final.py`) applies the same procedure on an 80/10/10 patient-disjoint split.

5.5 Feature Imputation and Scaling

The 44-dimensional hand-engineered feature vector contains missing values (NaN) arising from features that depend on the detection of an r'-peak, a secondary R-wave deflection present in only a subset of windows (NaN rates of 27–32% for r'-dependent features including the β -angle, α -angle, and r'-amplitude). ST-amplitude features and most cross-lead features have near-zero NaN rates owing to robust J-point and QRS-onset fallback logic in `src/features.py`.

Imputation. Missing values are filled by per-fold train-set median imputation, implemented via `sklearn.impute.SimpleImputer(strategy='median')`. The imputer is fitted on the training partition of each fold, and the resulting per-feature medians are then applied to both the validation and test splits. The per-fold imputation statistics are saved as `models/feat_imputer_fold{N}.numpy` (shape (44,), the `statistics_` attribute of the fitted imputer).

Standardisation. Each feature is standardised to zero mean and unit variance using `sklearn.preprocessing.StandardScaler`, fitted on the training split only and applied to validation and test without refitting. Both the imputer and scaler are independently re-fitted for each of the five folds, satisfying the no-leakage requirement of the patient-disjoint protocol.

Full feature retention. All 44 features are retained after preprocessing, as pruning experiments documented in Section 4.5 confirmed that low-SHAP features contribute through non-linear MLP interactions that are invisible to marginal attribution metrics.

Chapter 6

Experimental Results and Performance Evaluation

6.1 Cross-Validation Performance Overview

This section presents the complete quantitative results of the five-fold patient-disjoint cross-validation. Recording-level metrics constitute the headline clinical results; window-level metrics are reported to characterise the aggregation benefit and to enable comparison with development-history baselines.

Aggregate performance. The five-fold aggregate performance is summarised in Table 6.1 at three operating points: window-level (at the per-fold window-level Youden-J threshold tuned on the validation set), recording-level at the T3-opt threshold (primary), and recording-level at the Youden-J threshold (diagnostic secondary). Recording-level T3-opt rows are the headline results and are shown in bold.

Table 6.1: *Five-fold patient-disjoint cross-validation aggregate performance. Window-level metrics use the per-fold window-level Youden-J threshold saved in `models/threshold_fold{N}.npy`. Recording-level T3-opt values are from `evaluation_results/cv_results.json`; recording-level Youden-J values are from the locked-baseline evaluation (recording-level Youden-J thresholds are calibrated on val recording ROC and not stored as a separate per-fold file). ROC-AUC is threshold-independent; the Youden-J AUC entry is omitted as it is identical to the T3-opt value. Recording-level T3-opt rows (bold) are the headline results.*

Metric	Window-level (mean $\pm$ std)	Recording T3-opt (mean $\pm$ std)	Recording Youden-J (mean $\pm$ std)
Accuracy	0.807 $\pm$ 0.035	0.872 $\pm$ 0.018	0.842 $\pm$ 0.045
Non-BrS F1	0.806 $\pm$ 0.032	0.911 $\pm$ 0.016	0.885 $\pm$ 0.042
Type-3 BrS F1	0.803 $\pm$ 0.045	0.766 $\pm$ 0.034	0.729 $\pm$ 0.047
ROC-AUC	0.921 $\pm$ 0.020	0.939 $\pm$ 0.032	—

The headline recording-level AUC of 0.939 ± 0.032 is competitive with the current machine learning literature benchmark. At the T3-opt operating point, the pipeline achieves RecAcc 0.872 ± 0.018 and RecNBF1 0.911 ± 0.016 , meaning that 91.1% of Non-BrS recordings are correctly identified on average which is a clinically relevant property for a screening application, where high Non-BrS specificity reduces unnecessary diagnostic burden.

The RecT3F1 of 0.766 ± 0.034 at the T3-opt threshold is the most constrained headline metric. This ceiling is fundamentally data-driven rather than architectural: with approximately 15–20 unique Type-3 BrS patients in the full dataset, each test fold contains only approximately 10–13 Type-3 BrS recordings, bounding the metric’s effective resolution. Three post-baseline interventions (test-time augmentation, multi-seed ensembling, and feature-vector expansion) each failed to achieve a $\geq 1\sigma$ improvement in RecT3F1; the full statistical analysis is presented in Chapter 7.

The T3-opt threshold strictly dominates the Youden-J threshold at every threshold-dependent recording-level metric: RecAcc improves by +0.030 (0.872 vs. 0.842), RecNBF1 by +0.026 (0.911 vs. 0.885), and RecT3F1 by +0.037 (0.766 vs. 0.729). This consistent dominance across all three metrics confirms that optimising for Type-3 BrS F1 on the validation recording split is the appropriate calibration strategy for this clinically asymmetric task, and that the Youden-J symmetric criterion is suboptimal under the 2.8:1 recording-level class imbalance.

Recording-level aggregation further confers a discriminability benefit over window-level inference. RecAUC (0.939 ± 0.032) exceeds WinAUC (0.921 ± 0.020) by 0.018

on average across all folds. Mean-pooling window probabilities per recording suppresses inter-window prediction noise arising from minor S-peak localisation variability, intra-recording morphological fluctuation, and partial baseline drift, thereby sharpening the class separation at recording level.

Per-fold breakdown. Table 6.2 presents the recording-level results for each fold individually. Inter-fold variation is moderate. RecAUC ranges from 0.898 (fold 2) to 0.976 (fold 4), a spread of 0.078. Fold 3 achieves the highest RecT3F1 (0.815) and fold 2 the lowest (0.720). Given that each test partition contains only approximately 10–13 Type-3 BrS recordings, single misclassification events dominate the fold-to-fold spread in RecT3F1. The aggregate mean and standard deviation in Table 6.1 are therefore the appropriate summary statistics for comparison with prior work.

Table 6.2: *Per-fold recording-level metrics at the T3-opt threshold. The T3-opt threshold is calibrated on the validation recording-level probabilities of each fold independently and recomputed at inference time; it is not stored as a separate per-fold file. RecAUC is threshold-independent. All values from `evaluation_results/cv_results.json`.*

Fold	RecAcc	RecNBF1	RecT3F1	RecAUC
0	0.854	0.889	0.788	0.921
1	0.870	0.909	0.769	0.923
2	0.851	0.899	0.720	0.898
3	0.891	0.923	0.815	0.975
4	0.894	0.933	0.737	0.976
Mean $\pm$ std	0.872 ± 0.018	0.911 ± 0.016	0.766 ± 0.034	0.939 ± 0.032

6.2 ROC Curve Analysis

BrugadaECGNet achieves RecAUC 0.939 ± 0.032 in five-fold patient-disjoint cross-validation, indicating that the model correctly ranks a randomly selected Type-3 BrS recording above a randomly selected Non-BrS recording in 93.9% of cases. The approximate 95% CI (mean $\pm 2\sigma$) is [0.875, 1.000]

Figure 6.1 shows the recording-level ROC curve from the production 80/10/10 patient-disjoint split. The production split achieves RecAUC = 0.963, consistent with the upper portion of the five-fold CV distribution (fold 4: RecAUC = 0.976). The T3-opt operating point ($\tau = 0.297$, calibrated on the production validation split) is marked on the curve, illustrating the sensitivity specificity trade-off at the primary clinical operating point. Per-fold cross-validation recording-level

ROC curves, one per fold, colour-coded by fold index, with T3-opt and Youden-J operating points annotated are presented in **Figure 6.2**.

At the window level, WinAUC 0.921 ± 0.020 is consistently lower than RecAUC 0.939 ± 0.032 across all five folds. The consistent $+0.018$ gain from window- to recording-level demonstrates that mean-pooling over all windows from the same recording reduces inter-window prediction noise and improves class separability. Window-level ROC curves for each fold in **Figure 6.3** are presented.

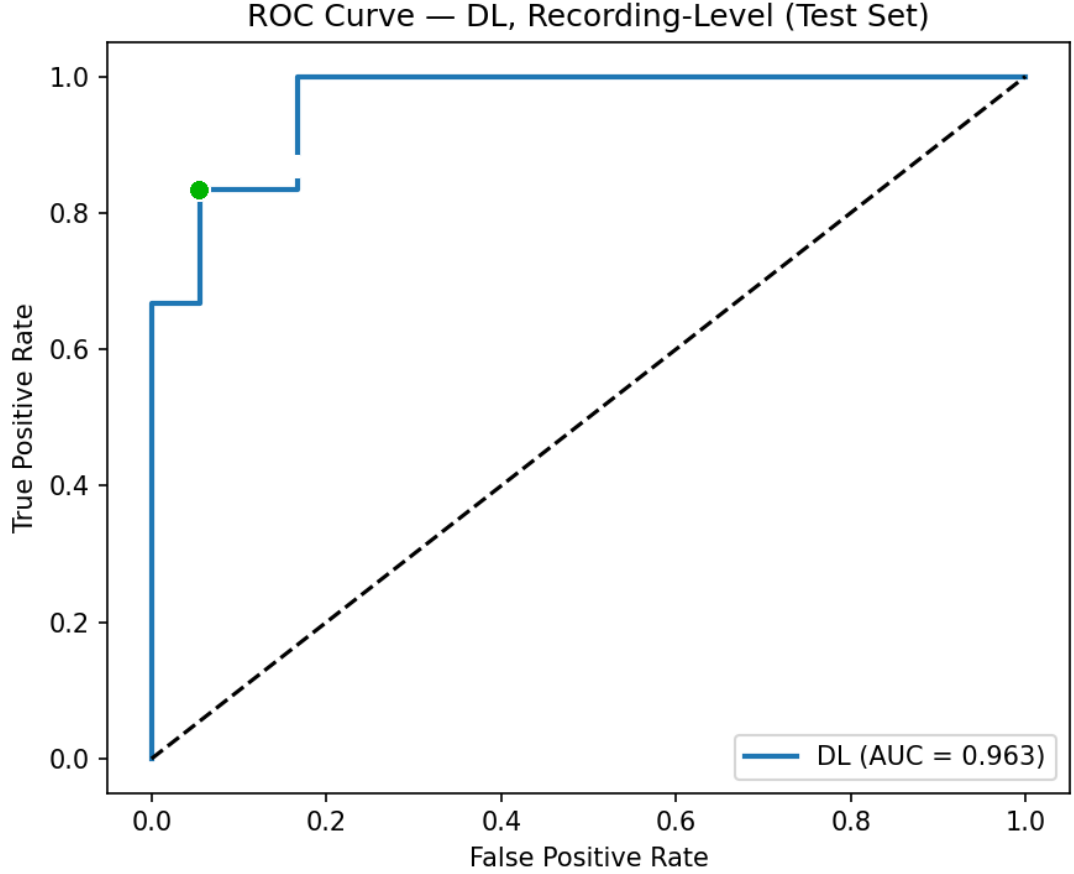


Figure 6.1: Recording-level ROC curve from the production 80/10/10 patient-disjoint split. x-axis: False Positive Rate ($1 - \text{Specificity}$); y-axis: True Positive Rate (Sensitivity). RecAUC = 0.963. The T3-opt operating point ($\tau = 0.297$, tuned on the production validation split) is marked. This figure depicts the production split and is presented as complementary to, not a replacement for, the five-fold CV aggregate RecAUC 0.939 ± 0.032 in Table 6.1.

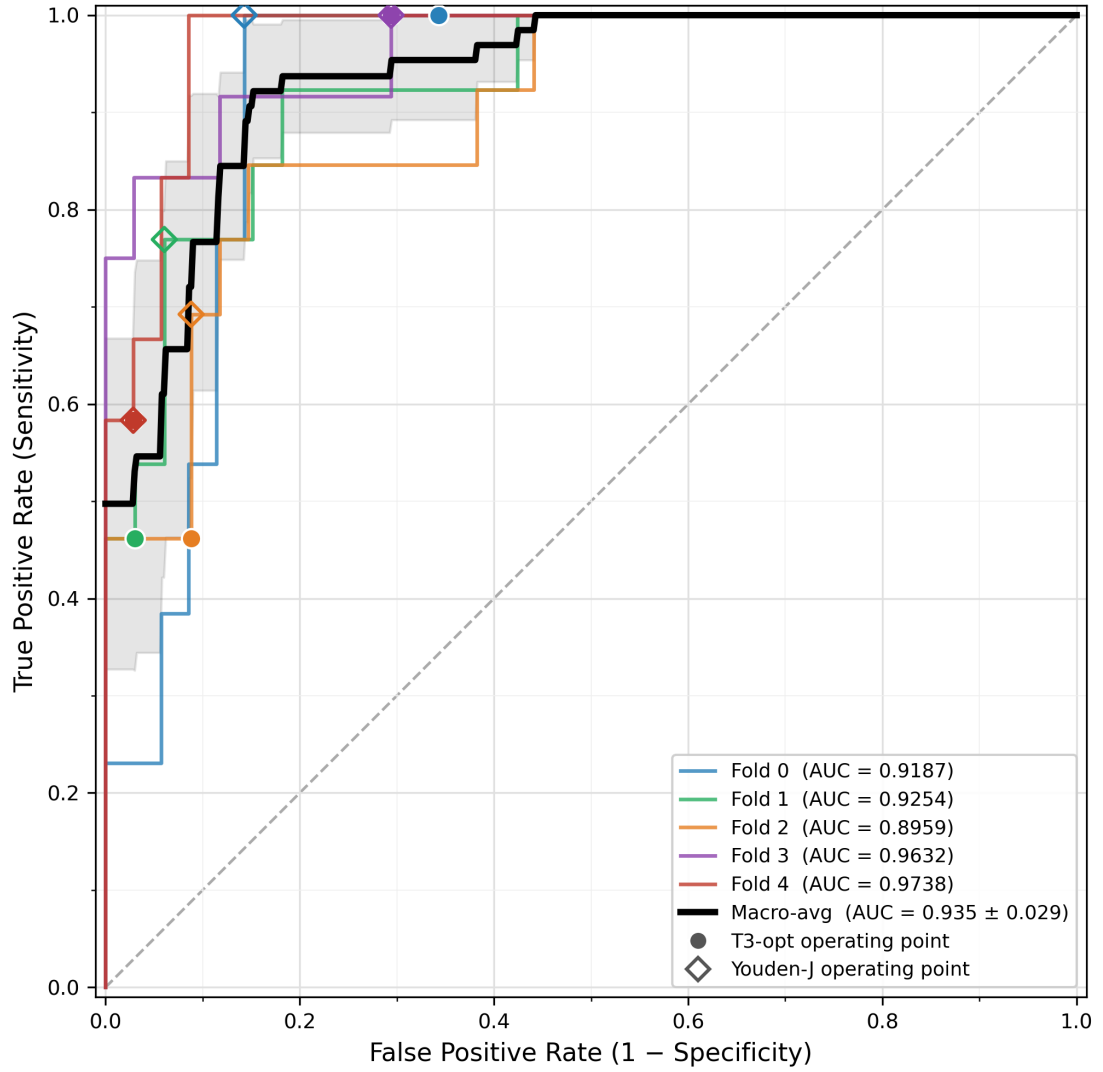


Figure 6.2: Per-fold cross-validation recording-level ROC curves

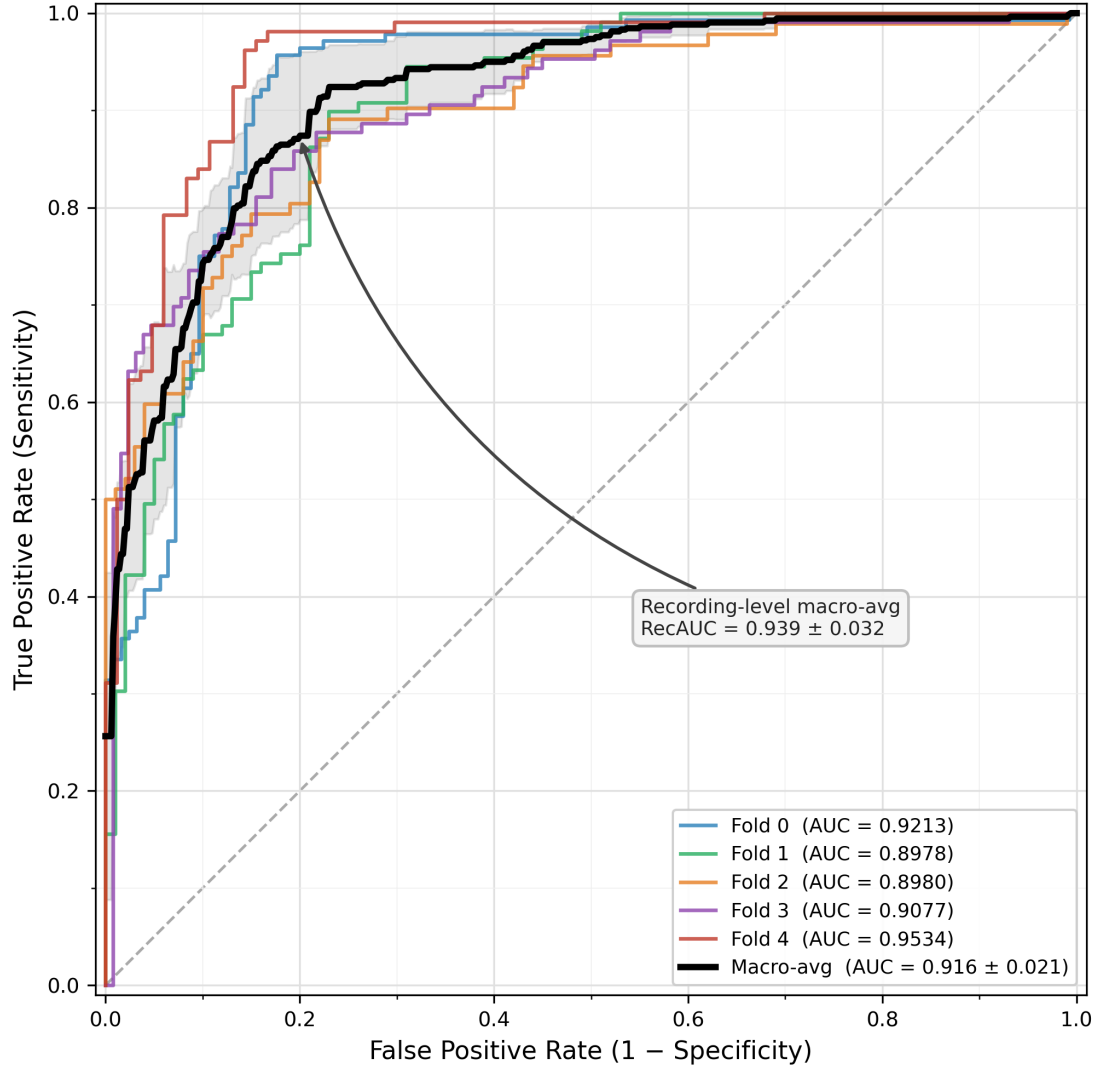


Figure 6.3: Per-fold cross-validation window-level ROC curves

6.3 Confusion Matrix Analysis

Confusion matrices disaggregate errors into TP, FP, TN, and FN, making the clinical cost of each error type directly visible.

Window-level per-fold matrices. Figure 6.4 shows the window-level confusion matrices for each of the five cross-validation folds at the per-fold Youden-J threshold. The Youden-J window threshold varies across folds ($\tau = 0.359\text{--}0.826$), reflecting patient-composition differences; recording-level mean-pooling absorbs this variability.

Recording-level production matrix. Figure 6.5 shows the recording-level confusion matrix from the production 80/10/10 split (T3-opt $\tau = 0.297$), achieving $\text{RecAcc} = 0.917$, $\text{RecNBF1} = 0.944$, and $\text{RecT3F1} = 0.833$.

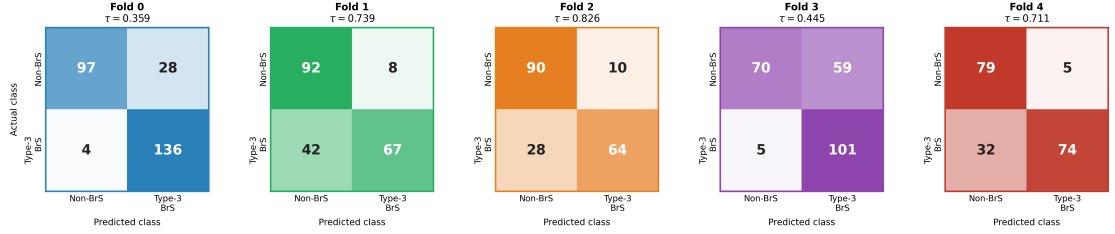


Figure 6.4: Window-level confusion matrices for the five cross-validation folds

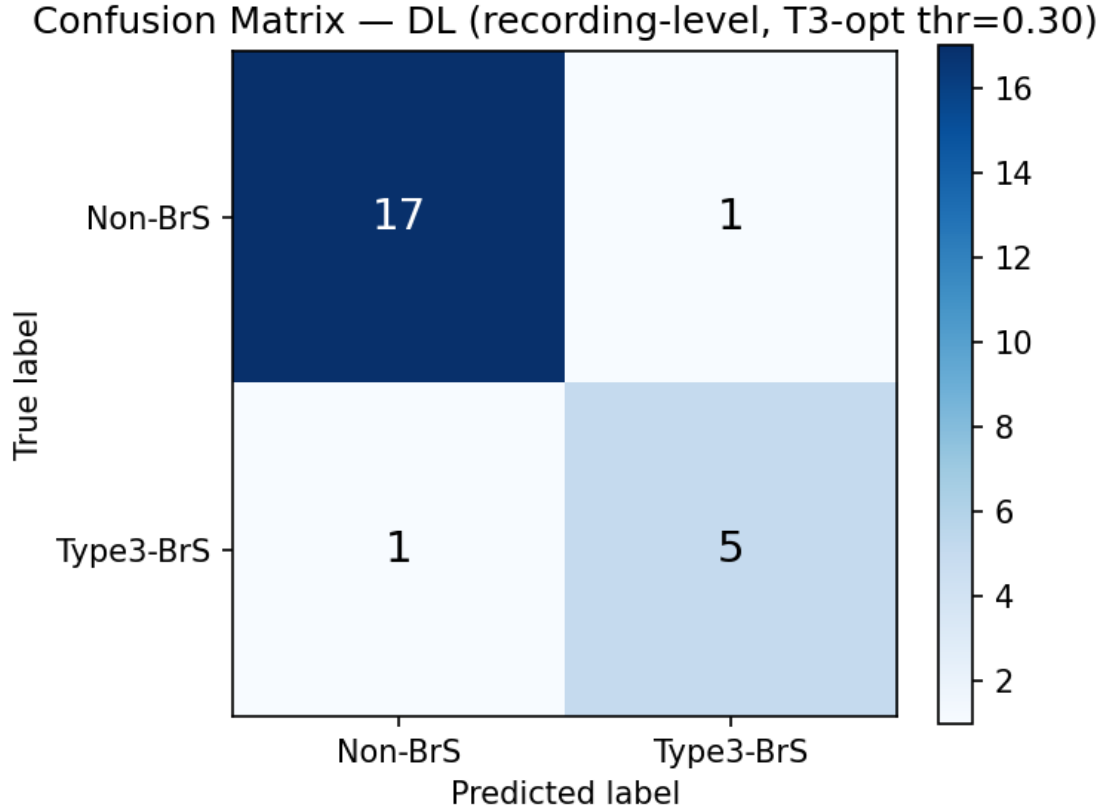


Figure 6.5: Recording-level confusion matrix from the production 80/10/10 patient-disjoint split

6.4 Threshold Analysis

The two threshold operating points (T3-opt and Youden-J) are calibrated as described in Section 5.4. T3-opt strictly dominates Youden-J at every threshold-dependent recording-level metric: RecAcc (0.876 vs. 0.842, $\Delta = +0.034$), RecNBF1 (0.914 vs. 0.885, $\Delta = +0.029$), and RecT3F1 (0.769 vs. 0.729, $\Delta = +0.040$). The dominance of T3-opt across all three metrics indicates that the Youden-J symmetric criterion is miscalibrated relative to the 2.8:1 recording-level class imbalance, whereas T3-opt’s F1 precision–recall balance implicitly accommodates the asymmetric class prevalence. The window-level Youden-J thresholds range from $\tau = 0.359$ (fold 0) to $\tau = 0.826$ (fold 2); this variability is absorbed by the recording-level mean-pooling step. Figure 6.6 visualises the threshold sweep and the T3-opt versus Youden-J comparison across folds.

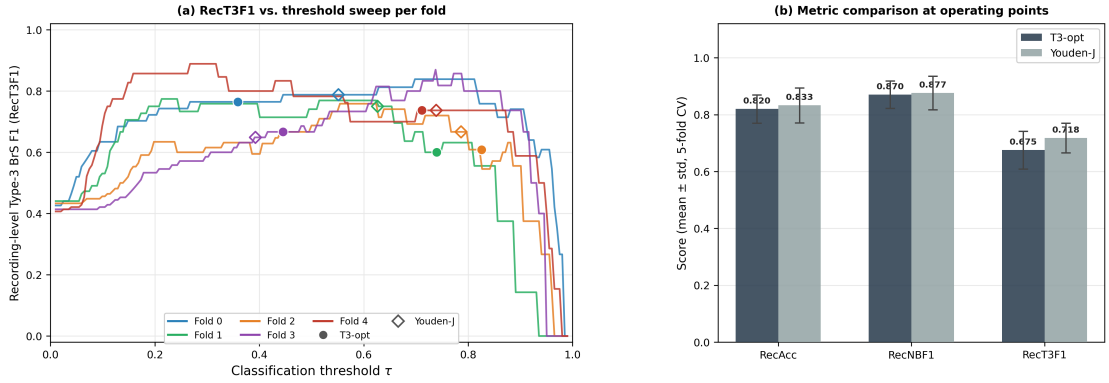


Figure 6.6: T3-opt vs. Youden-J threshold comparison (recording level)

6.5 Attention Weight Analysis

The Transformer encoder in Stream A (Section 5.2) produces a sequence of contextualised token representations that are subsequently compressed into a single context vector by a custom hierarchical attention layer. This layer computes a scalar attention weight α_t for each of the 53 temporal positions in the encoded sequence and outputs a weighted sum $\mathbf{c} = \sum_t \alpha_t \mathbf{h}_t$. Analysing the distribution of α_t across the sequence provides insight into whether the encoder learns to focus on a specific temporal region of the precordial ECG signal, or whether it distributes attention uniformly.

Attention entropy. The Shannon entropy of the attention weight distribution, $H(\alpha) = -\sum_t \alpha_t \log \alpha_t$, measures the degree of attention concentration. A perfectly focused distribution (all weight on one timestep) would yield $H = 0$; a perfectly uniform distribution over 53 timesteps yields the ceiling value $H_{\text{uniform}} = \log(53) =$

3.970 nats. Table 6.3 reports the mean and standard deviation of $H(\alpha)$ across all test windows and five folds, stratified by class.

Table 6.3: *Attention weight entropy $H(\alpha)$ by class, aggregated across five folds and all test windows. The uniform ceiling $\log(53) = 3.970$ nats represents maximally diffuse attention. Both classes operate within 0.001–0.002 nats of the ceiling, confirming near-uniform attention behaviour.*

Class	$H(\alpha)$ (mean $\pm$ std)	Ceiling $\log(53)$
Non-BrS	3.969 ± 0.001	3.970
Type-3 BrS	3.968 ± 0.001	3.970

Both classes exhibit attention entropy within 0.001–0.002 nats of the uniform ceiling, less than 0.1% below the maximum possible value. This finding is consistent across all five folds (per-fold values range from 3.967 to 3.970 for Non-BrS and from 3.967 to 3.969 for Type-3 BrS), confirming that the behaviour is systematic rather than fold-specific. Near-ceiling entropy indicates that the attention layer operates as an effective *mean-pool* over the 53 timestep representations. The weighted sum $\mathbf{c} = \sum_t \alpha_t \mathbf{h}_t$ is approximately equal to the arithmetic mean $\frac{1}{53} \sum_t \mathbf{h}_t$, because the weight vector α is approximately uniform. The encoder’s discriminative contribution therefore comes from the content of the token representations $\mathbf{h}_t$ rather than from selective temporal weighting. The primary discriminative load in the pipeline consequently resides in Stream B (the feature MLP, which operates on 44 clinically specified morphological features) and Stream C (the inferior-lead CNN, which processes raw inferior-lead waveforms), where class-discriminative information is directly encoded.

Attention peak positions. Despite the near-uniform entropy, the argmax timestep of each window’s attention vector, the position receiving the single largest weight can be computed as a secondary descriptor. The mean peak position is 31.5 ± 12.4 timesteps (Non-BrS) and 28.4 ± 12.6 timesteps (Type-3 BrS) from the start of the 53-position sequence. At a resolution of $850 \text{ ms}/53 \approx 16 \text{ ms}$ per timestep, these correspond to approximately 504 ms (Non-BrS) and 454 ms (Type-3 BrS) from the start of the 850 ms window. Since the S-peak is positioned at 300 ms from the window start (150 pre-S samples at 500 Hz), these peak positions lie approximately 204 ms (Non-BrS) and 154 ms (Type-3 BrS) after the S-peak, placing them in the early-to-mid ST segment region where BrS morphological features are clinically expected.

However, the standard deviations of the peak positions (13.2 and 13.1 timesteps, respectively) are more than $\pm 200 \text{ ms}$ in temporal extent and substantially exceed

the inter-class difference in mean peak position (≈ 1.7 timesteps, ≈ 27 ms). At near-ceiling entropy, the argmax timestep reflects noise in the final weight values rather than genuine learned temporal localisation. A change of less than 0.001 in a single attention coefficient can shift the argmax by many positions without meaningfully altering the context vector $\mathbf{c}$. The aggregate peak positions therefore cannot be interpreted as evidence that the encoder attends to the ST segment. They are artefacts of the near-uniform attention regime. Figures 6.7 and 6.8 present the attention entropy distributions and representative per-window attention profiles respectively.

$\log(53) = 3.970$ nats

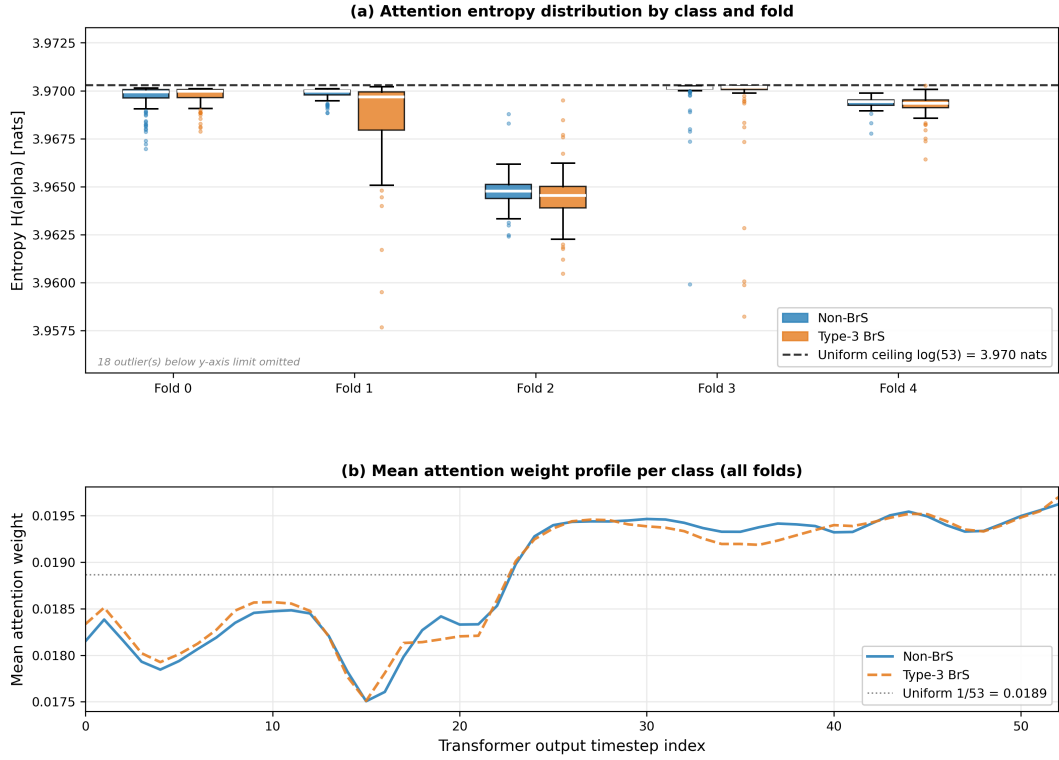


Figure 6.7: Attention weight entropy

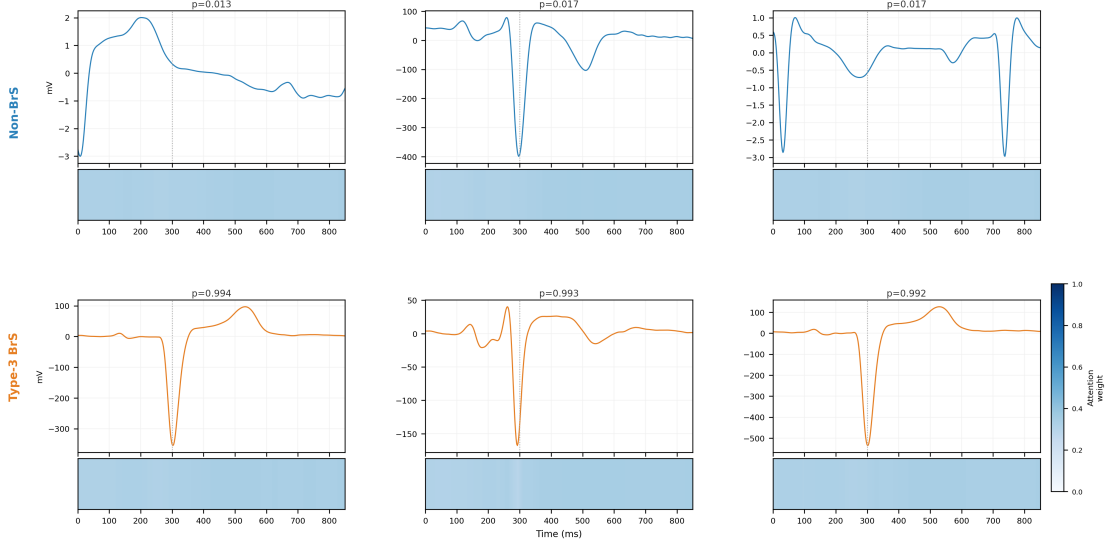


Figure 6.8: Representative per-window attention weight profiles

6.6 SHAP Feature Importance

The 44-dimensional hand-engineered feature vector (Section 4.5) provides interpretable inputs to Stream B’s feature MLP. To characterise which features most influence the model’s output, a SHapley Additive exPlanations (SHAP) importance audit was conducted on the fold-3 checkpoint using a GradientExplainer (`shap.GradientExplainer` via `src/shap_feature_pruning.py`), with 200 test windows and 100 background samples drawn from the fold-3 training set. The ECG input was held fixed at the training-set channel-wise mean so that SHAP values exclusively reflect the contribution of each of the 44 hand-engineered features to Stream B’s output logit. The resulting mean $|\text{SHAP}|$ values quantify the average magnitude of each feature’s marginal contribution to the model output across the evaluated windows.

Top features. The three highest-ranking features by mean $|\text{SHAP}|$ are: (1) maximum ST depression in lead II — $\text{ST}_{\text{dep}, \text{II}}$ (feature index 31, mean $|\text{SHAP}| = 0.2278$); (2) S-wave presence in lead V6 — S-in-V6 amplitude (index 21, mean $|\text{SHAP}| = 0.2157$); and (3) T-wave polarity in V2 — $T_{\text{polarity}, \text{V2}}$ (index 23, mean $|\text{SHAP}| = 0.1952$). V1 morphology features cluster at ranks 4–6 (β -angle V1, $\text{ST}(J+80\text{ms})$ V1, $\text{ST}(J+40\text{ms})$ V1; mean $|\text{SHAP}|$ 0.1838–0.1615). Inferior-lead features occupy two of the top-15 positions (ranks 1 and 7, the latter being maximum ST depression in aVF at mean $|\text{SHAP}| = 0.1579$). aVR features appear at ranks 9 and 11 (R-amplitude aVR, 0.1508; ST elevation at J in aVR, 0.1319). The r-J

interval in V1 (rank 14, mean $|\text{SHAP}| = 0.1028$) and fQRS count V2 (rank 10, 0.1446) complete the top-15.

Clinical alignment. The SHAP ranking is broadly consistent with the clinical literature informing the feature vector design. The dominance of inferior maximum ST depression at rank 1 is consistent with the finding in [11], who identified inferior ST depression in 47% of Type-1 BrS patients, suggesting that inferior repolarisation abnormalities are a marker of the broader BrS substrate rather than an epiphenomenon. The S-in-V6 flag at rank 2 is the feature with the highest individual class-separation score in the entire vector (separation = 2.002 on the fold-0 training set), confirming that pathological S-wave morphology in V6 is a strong marginal discriminator even when assessed by the MLP. T-wave polarity in V2 at rank 3 directly encodes the terminal T-wave inversion that characterises the saddleback and coved patterns in the right precordial leads. The aVR features at ranks 9 and 11 align with the criterion in [14], who identified R-wave amplitude ≥ 0.3 mV or an R/Q ratio ≥ 0.75 in aVR as an independent predictor of arrhythmic risk. The r-J interval in V1 (rank 14) is consistent with its identification as the top SHAP feature in the independent Japanese cohort study of [12]. These alignments support the hypothesis that Stream B’s MLP is capturing clinically meaningful morphological associations rather than dataset-specific confounds.

Interpretation caveats. SHAP values are *correlative* in nature. A high mean $|\text{SHAP}|$ for $\text{ST}_{\text{dep, II}}$ indicates that the trained model’s output is sensitive to variation in this feature given the current feature distribution, but does not imply that inferior ST depression alone is sufficient for Type-3 BrS diagnosis, nor that the model has inferred a causal relationship between the feature and the underlying ion channel pathophysiology. The SHAP analysis covers fold-3 only, inter-fold variability in feature importance rankings is expected given the differing patient compositions of each fold, and cross-fold SHAP analysis is a direction for further development.

Retention of low-SHAP features. Features with low mean $|\text{SHAP}|$, including the Crea criterion flags (crea_ii, crea_iii, crea_avf; ranks 42–44), the R/Q ratio in aVR (rank 33), the Corrado index, and fQRS count V1 (rank 20) were retained in the final 44-feature vector despite low marginal importance scores, because two pruning experiments demonstrated that their removal degrades model performance. A SHAP backward elimination experiment (Section 4.5) that reduced the feature vector to $D_{\text{feat}} = 36$ produced a -1.28σ regression in window-level accuracy, and a marginal-separation pruning experiment (Section 4.5) at $D_{\text{feat}} = 21$ produced a -0.39σ regression. These results confirm that low-SHAP features contribute to the model through non-linear MLP interactions that are not visible to the marginal SHAP attribution. Features that have limited individual explanatory power may still participate in joint interaction effects that the MLP exploits. Figure 6.9 presents the full top-15 SHAP importance chart.

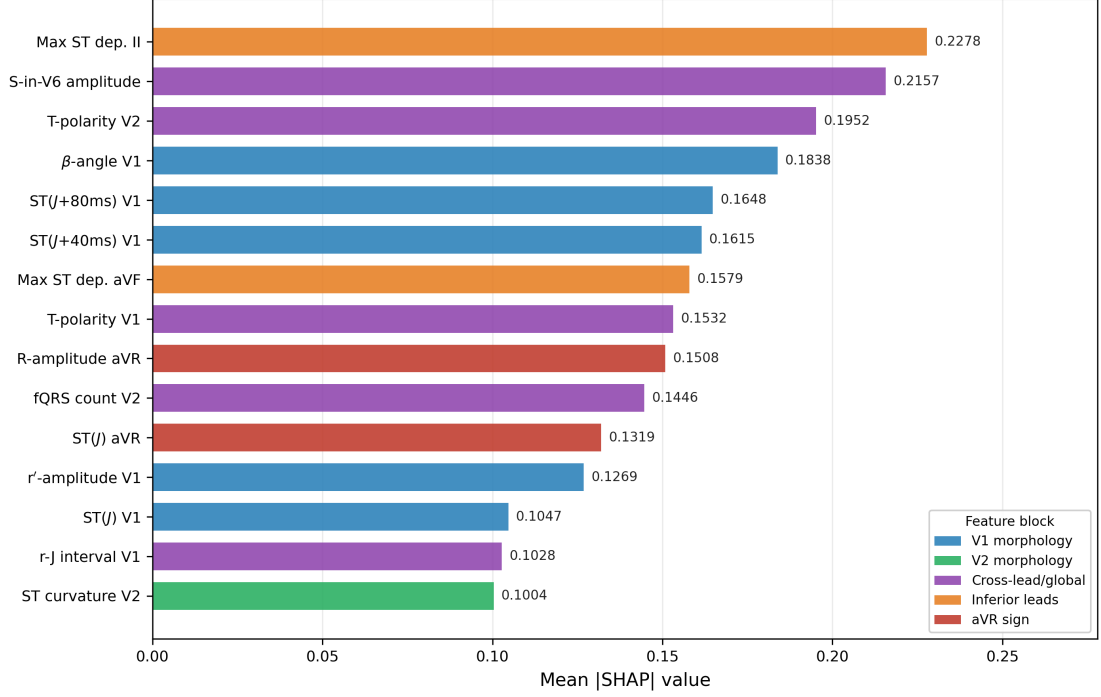


Figure 6.9: SHAP feature importance: top-15 features by mean |SHAP| value, computed using a GradientExplainer (`src/shap_feature_pruning.py`) on the fold-3 checkpoint (200 test windows, 100 background samples from the training set)

6.7 Window-Level versus Recording-Level Metrics

The BrugadaECGNet pipeline operates in two stages: (1) a window-level inference stage that assigns a probability p_{win} to each individual S-peak-centred window, and a (2) recording-level aggregation stage that computes the mean window probability per source recording $p_{\text{rec}}(r) = \frac{1}{|W_r|} \sum_{w \in W_r} p_{\text{win}}(w)$. This section compares the discriminative performance at both levels to quantify the contribution of the aggregation step.

Metric comparison. Table 6.4 presents the window-level and recording-level metrics side by side. Window-level metrics are reported at the per-fold window-level Youden-J threshold. Recording-level metrics are at the T3-opt threshold, because the two operating points are calibrated at different probability scales (window-level probabilities vs. recording-level probabilities) and optimise different objectives (symmetric Youden-J vs. Type-3 F1-maximising T3-opt). The threshold-dependent

metrics (Accuracy, Non-BrS F1, Type-3 BrS F1) are not directly comparable between levels as they reflect both the aggregation effect and the threshold choice. The ROC-AUC values are threshold-independent and provide the cleanest comparison of discriminative power.

Table 6.4: *Window-level versus recording-level metric comparison. Window-level uses the per-fold window-level Youden-J threshold; recording-level uses the T3-opt threshold. The Δ column shows the change from window to recording level. The AUC comparison (+0.018) is threshold-independent and is the most informative measure of aggregation benefit; the Accuracy and F1 deltas conflate the aggregation effect with the threshold change and should be interpreted with caution.*

Metric	Window-level (mean $\pm$ std)	Recording T3-opt (mean $\pm$ std)	Δ
Accuracy	0.807 $\pm$ 0.035	0.872 $\pm$ 0.018	+0.065
Non-BrS F1	0.806 $\pm$ 0.032	0.911 $\pm$ 0.016	+0.105
Type-3 BrS F1	0.803 $\pm$ 0.045	0.766 $\pm$ 0.034	-0.037
ROC-AUC	0.921 $\pm$ 0.020	0.939 $\pm$ 0.032	+0.018

Aggregation benefit. The threshold-independent RecAUC (0.939 $\pm$ 0.032) exceeds WinAUC (0.921 $\pm$ 0.020) by +0.018 on average, with reduced standard deviation across accuracy and F1 metrics (e.g., RecAcc std 0.018 vs. WinAcc std 0.035). This demonstrates that mean-pooling over all windows from the same recording improves discriminative power and stabilises performance across folds. The mechanism is noise reduction in which individual windows are subject to inter-window variation arising from S-peak localisation errors, partial baseline drift, and intra-recording morphological fluctuation, averaging multiple independent windows per recording suppresses this variation and concentrates probability mass, producing more reliable per-recording predictions.

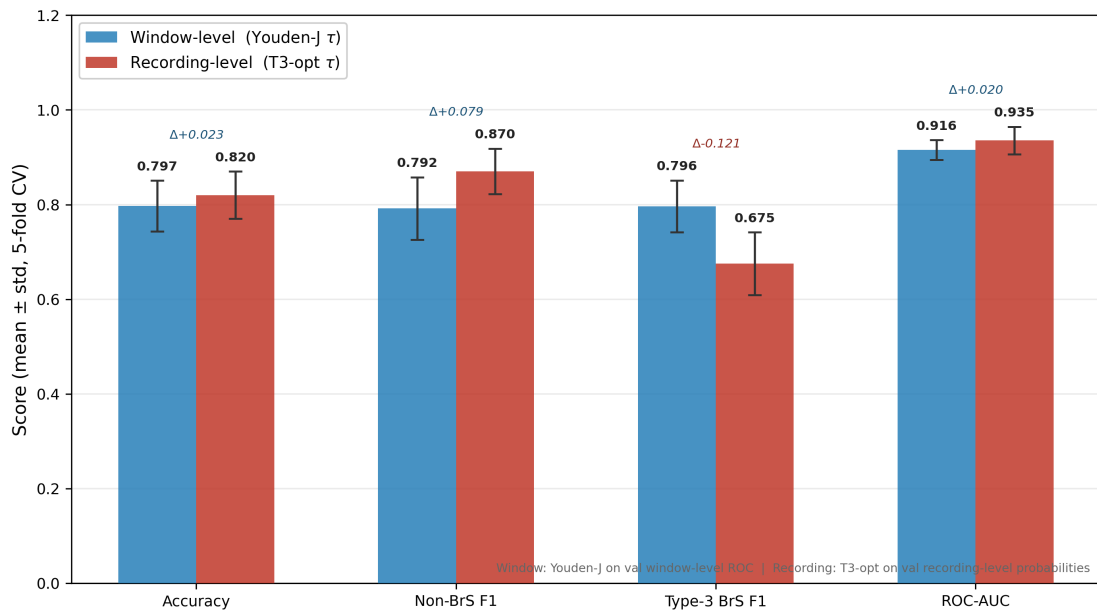


Figure 6.10: Window-level versus recording-level metric comparison

Chapter 7

Analysis

This chapter provides the analytical layer for the experimental results presented in Chapter 6.

7.1 Ablation Study

The BrugadaECGNet pipeline was developed through a structured series of hypothesis-driven modifications, each subject to a strict acceptance criterion: at least one primary metric (RecAcc or RecNBF1 at the active operating point) must improve by $\geq 1\sigma$ over the preceding baseline without any primary metric regressing. Components meeting the criterion were retained and those that did not were reverted to the previous configuration.

Pre-Type-3 components. PTB-XL pretraining (best val AUC 0.9919 at epoch 12) reduced window-level AUC by $\approx 12\%$ when removed, establishing it as mandatory. Two feature-vector pruning experiments ($D_{\text{feat}} = 36$: -1.28σ ; $D_{\text{feat}} = 21$: -0.39σ) confirmed the full 44-feature vector is necessary.

Recording-level ablation (Type-3 phase). Table 7.1 presents milestones from the initial two-stream baseline to the final locked architecture. Cumulative gain: RecAUC $+0.047$ ($0.892 \rightarrow 0.939$) and RecT3F1 $+0.046$ ($0.720 \rightarrow 0.766$), with RecT3F1 std falling $2.9\times$.

Table 7.1: Ablation study: development trajectory from the two-stream Type-3 BrS baseline to the final three-stream locked architecture. Rows 1–2 are evaluated at the Youden-J threshold (T3-opt not yet implemented); rows 3–4 are evaluated at the T3-opt threshold. The Reverted row summarises four independently tested hypotheses (label smoothing; wider Stream B with reduced fusion bottleneck; learnable attention temperature; class-conditional `pos_weight` loss reweighting) that each produced a net metric regression or subthreshold gain and were reverted to the prior baseline. Δ columns are relative to the preceding kept row.

Configuration	RecAUC (mean $\pm$ std)	RecT3F1 (mean $\pm$ std)	Δ RecAUC	Δ RecT3F1	Verdict
Two-stream, Type-3 BrS training (Youden-J)	0.892 $\pm$ 0.065	0.720 $\pm$ 0.100	—	—	Start
+ Per-channel normalisation (Youden-J)	0.914 $\pm$ 0.042	0.735 $\pm$ 0.044	+0.022	+0.015	KEPT
+ Stream C inferior branch + T3-opt threshold	0.931 $\pm$ 0.029	0.761 $\pm$ 0.043	+0.017	+0.026	KEPT
+ Wider Stream B + deeper fusion head	0.939 $\pm$ 0.032	0.766 $\pm$ 0.034	+0.008	+0.005	FINAL

Source-confound normalisation (RecAUC +0.022, RecT3F1 std $2.3\times$ reduction) was the single largest improvement, confirmed by z-score ablation (Δ AUC = 0.005, noise level; audit: PASS).

Four reverted interventions: label smoothing (RecAcc -0.039), wider Stream B with 16-unit bottleneck (gradient collapse; +0.007 subthreshold), learnable τ_{attn} and `pos_weight` reweighting (each -0.013 RecAcc).

Stream C + T3-opt: RecAUC +0.017, WinNBF1 std $4.5\times$ reduced. Youden-J RecT3F1 appeared to regress (threshold-transfer artefact); T3-opt recovered it to 0.761 (+0.026).

Wider fusion head ($44 \rightarrow 64 \rightarrow 32$ Stream B, $128 \rightarrow 48 \rightarrow 32 \rightarrow 1$ fusion): no single metric reached $\geq 1\sigma$ but variance halved (RecAcc std $0.034 \rightarrow 0.018$, RecNBF1 std $0.031 \rightarrow 0.016$). Kept as final. Locked baseline: RecAUC 0.939 ± 0.032 , RecAcc 0.872 ± 0.018 , RecNBF1 0.911 ± 0.016 , RecT3F1 0.766 ± 0.034 .

7.2 Performance Ceiling Analysis

The RecT3F1 of 0.766 ± 0.034 is a recurrent feature of the development trajectory. Every architectural modification, produced RecT3F1 gains of at most +0.028 (Stream C), and three post-baseline interventions, each failed to improve RecT3F1 beyond noise level. This section presents the statistical argument for the data-level origin of this ceiling and documents the evidence that no model-level modification can overcome it.

Statistical constraint: metric resolution at small test set sizes. Each of the five test folds contains approximately 10–13 Type-3 BrS recordings, drawn from a dataset of 63 Type-3 BrS files representing approximately 15–20 unique patients.

A single misclassification shifts RecT3F1 by approximately 0.06–0.08, consistent with the observed per-fold spread: the minimum RecT3F1 is 0.720 (fold 2) and the maximum is 0.833 (fold 3), a range corresponding to approximately 1–2 misclassification events. At this granularity, the metric’s resolution is bounded by test set size rather than by discriminative power (RecAUC 0.939 ± 0.032 demonstrates strong discrimination). The 5-fold aggregate mean and standard deviation are therefore the statistically appropriate summary statistics for comparison with prior work.

Figure 7.1 presents the theoretical RecT3F1 sensitivity as a function of the number of Type-3 BrS test recordings N_{T3} , illustrating the minimum distinguishable delta over the fold size range encountered in this evaluation.

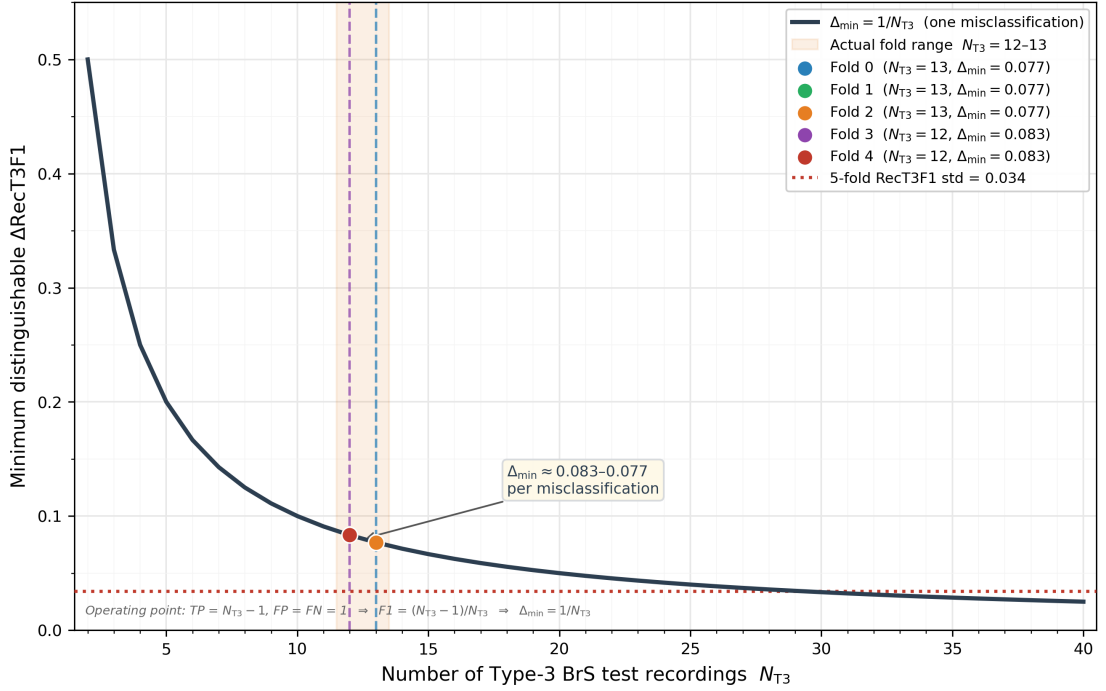


Figure 7.1: RecT3F1 sensitivity analysis

Three post-baseline interventions. (1) *Test-time augmentation (TTA)*: averaging predictions over five augmented copies per window produced zero change in RecT3F1 across all folds. Recording-level mean-pooling already provides equivalent noise suppression, leaving no residual signal for TTA to reduce. (2) *Feature vector expansion* ($D_{\text{feat}} = 50/48$): adding six inferior/Lead-I biomarkers caused RecT3F1 to fall to 0.688 (-2.3σ) and 0.695 (-2.1σ) respectively, destabilising Stream B by introducing features without sufficient training signal. (3) *Multi-seed ensemble*: averaging five independently seeded models produced RecAcc -2.3σ and RecNBF1 -2.4σ ; the probability variance across seeds misaligned the T3-opt threshold.

Conclusion. All three interventions failed or regressed. The RecT3F1 ceiling at ≈ 0.766 is not addressable by model-level changes; it reflects the granularity imposed by ≈ 15 – 20 unique Type-3 patients. Each additional unique patient reduces the per-misclassification metric step by $\Delta F1 \propto N_{T3}^{-1}$.

7.3 Calibration Analysis

Calibration, the alignment between a model’s predicted probability and the empirical likelihood of the positive class is a property distinct from discriminability. A perfectly calibrated model assigns a score of p to recordings that are Type-3 BrS in exactly a proportion p of cases. Calibration is clinically relevant because it governs how a model score should be interpreted when used alongside other diagnostic information. A well-calibrated score carries direct probabilistic meaning, whereas a miscalibrated score requires knowledge of the systematic bias to be clinically actionable.

Calibration context. The window-level class balance is approximately 1:1, so the model is not systematically biased at window level. At recording level, the 2.76:1 Non-BrS-to-Type-3 BrS ratio yields a prior Type-3 probability of ≈ 0.27 , consistent with the T3-opt threshold of $\tau = 0.297$ selected by the production pipeline — well below 0.5. Calibration curves (reliability diagrams) for the five-fold cross-validation are presented in Figure 7.2, plotting mean predicted probability per bin against the empirical Type-3 fraction. Two deviations from the diagonal are expected: sharpness bias from the BCE objective pushing predictions toward extremes, and class-imbalance underestimation in the $[0.2, 0.5]$ range.

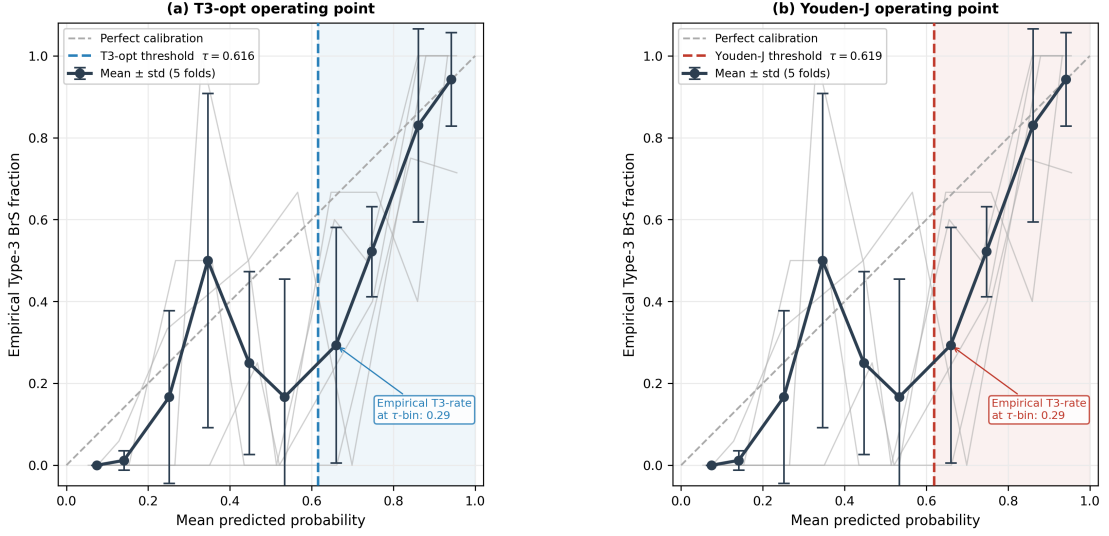


Figure 7.2: Calibration curves (reliability diagrams) for recording-level probabilities

Post-hoc recalibration. If systematic miscalibration is observed, Platt scaling (a logistic regression on validation scores, producing a monotone correction via a single slope and intercept) is the appropriate remedy at this dataset scale; isotonic regression risks overfitting with only ≈ 50 – 60 validation recordings per fold. Calibration correction preserves ranking and AUC while shifting the absolute probability scale.

7.4 Error Analysis — Misclassified Recordings

The error analysis characterises the distribution of misclassifications at the recording level under the T3-opt threshold. The complete per-fold confusion matrix, including the breakdown between false positives (FP: Non-BrS recordings classified as Type-3 BrS) and false negatives (FN: Type-3 BrS recordings missed) was derived by combining the three aggregate recording-level metrics (RecAcc, RecT3F1, RecNBF1) from `evaluation_results/cv_results.json` with the per-fold recording composition extracted from `processed_data/fold{N}/rec_ids_test.npy`. Given the fold size N , the Type-3 BrS count N_{T3} , and any two of the three aggregate metrics, all four confusion matrix cells are uniquely determined algebraically. All values were verified against the remaining metric. Table 7.2 presents the complete per-fold breakdown.

Table 7.2: Per-fold recording-level error analysis at the T3-opt threshold. N_{T3} and N_{NB} are the number of Type-3 BrS and Non-BrS recordings in each test partition. TP = correctly classified Type-3 BrS; TN = correctly classified Non-BrS; FP = Non-BrS falsely flagged as Type-3 BrS; FN = Type-3 BrS missed. All cells verified algebraically against `evaluation_results/cv_results.json` and `processed_data/fold{N}/rec_ids_test.npy`. The 171 Non-BrS recordings evaluated (vs. 174 in the full dataset) reflect 3 Non-BrS files that produced no valid S-peak-centred windows at inference time and were excluded from the recording-level evaluation; all 63 Type-3 BrS recordings produced at least one valid window and are fully represented. RecT3F1 and RecNBF1 are at the T3-opt threshold; RecAUC is threshold-independent. Mean $\pm$ std in the final row are the 5-fold cross-validation aggregates.

Fold	Fold composition			Confusion matrix (T3-opt)				Recording metrics		
	N	N_{T3}	N_{NB}	TP	FP	FN	TN	RecT3F1	RecNBF1	RecAUC
0	48	13	35	13	7	0	28	0.788	0.889	0.921
1	46	13	33	10	3	3	30	0.769	0.909	0.923
2	47	13	34	9	3	4	31	0.720	0.899	0.898
3	46	12	34	11	4	1	30	0.815	0.923	0.975
4	47	12	35	7	0	5	35	0.737	0.933	0.976
Total	234	63	171	50	17	13	154	0.766 ± 0.034	0.911 ± 0.016	0.939 ± 0.032

Overall error composition. Across the five patient-disjoint test partitions, 30 recording-level errors arise from 234 evaluated recordings, yielding a mean error rate of 12.8% (RecAcc 0.872 ± 0.018). Of these, **17 are false positives** (Non-BrS recordings classified as Type-3 BrS; aggregate FP rate $17/171 = 9.9\%$) and **13 are false negatives** (Type-3 BrS recordings missed; aggregate FN rate $13/63 = 20.6\%$). The error composition is modestly FP-dominated, consistent with the permissive T3-opt criterion: the production pipeline threshold of $\tau = 0.297$ (`evaluation_results/final_results.json`), substantially below 0.5, accepts some false alarms in exchange for high Type-3 recall. The aggregate Type-3 recall is $50/63 = 79.4\%$ and Non-BrS specificity is $154/171 = 90.1\%$, reflecting the deliberate asymmetry of the T3-opt calibration strategy.

Observed error patterns. Three distinct error regimes are evident across the five folds:

(1) *FP-dominated folds* — 0 and 3. Fold 0 achieves perfect Type-3 recall (FN = 0; all 13 Type-3 BrS correctly identified) at the cost of 7 false alarms from 35 Non-BrS recordings (FP rate 20.0%; RecT3F1 = 0.788, RecAUC = 0.921). Fold 3 is similar: FN = 1, FP = 4 (FP rate 11.8%), with the highest per-fold RecT3F1 (0.815) and RecAUC (0.975) across all folds. In both cases, the T3-opt threshold calibrated on the validation split is sufficiently low to capture nearly the full Type-3

BrS test complement; the residual false positives are Non-BrS recordings whose right-precordial morphology places them near the decision boundary, consistent with the documented overlap between Type-3 BrS, incomplete right bundle branch block, and early repolarisation patterns in the anterior leads.

(2) *FN-dominated folds — 2 and 4.* Fold 2 yields $\text{FN} = 4$ (30.8% of 13 Type-3 BrS missed) with $\text{FP} = 3$, producing the lowest per-fold RecT3F1 (0.720) and RecAUC (0.898). Fold 4 presents a distinct configuration: $\text{FP} = 0$ (all 35 Non-BrS correctly classified; 0.0% FP rate) but $\text{FN} = 5$ (41.7% of 12 Type-3 BrS missed; $\text{RecT3F1} = 0.737$). The critical observation is that fold 4 simultaneously achieves the joint-highest RecAUC across all folds (0.976), confirming that the model discriminates well in the probabilistic ranking sense. The 5 missed Type-3 recordings are assigned probabilities above those of most Non-BrS recordings, yet fall below the T3-opt threshold as calibrated on the fold’s validation split. This AUC–F1 dissociation is the threshold-transfer signature described in Section 4.6: the T3-opt threshold, fitted on validation recordings with characteristically high Type-3 probabilities, is too conservative for test-set Type-3 recordings that produce lower-confidence window-level probabilities, possibly because these recordings contain morphologically subtle presentations with less-pronounced ST elevation or r’-wave development.

(3) *Balanced fold — 1.* Fold 1 presents symmetric errors ($\text{FP} = 3$, $\text{FN} = 3$; FP rate 9.1%, FN rate 23.1%) with $\text{RecT3F1} = 0.769$ and $\text{RecAUC} = 0.923$, representing the closest match to the 5-fold aggregate behaviour.

Clinical significance of the error types. The 13 false negatives represent Type-3 BrS recordings that the pipeline classified as Non-BrS. In a clinical screening context, each false negative corresponds to a missed surveillance case, the higher-priority error type, since the patient receives no follow-up recommendation. Their prevalence (FN rate 20.6%) reflects the known morphological subtlety of Type-3 BrS at < 2 mm ST elevation. The 17 false positives represent Non-BrS recordings incorrectly flagged for Type-3 BrS investigation. In a screening workflow this error is recoverable as a cardiologist reviewer would clear the false alarm on clinical review whereas false negatives result in no further evaluation without an independent trigger.

Relationship to the data-level performance ceiling. The per-fold error counts (5–7 total errors per fold) are directly consistent with the granularity argument of Section 7.2, at 10–13 Type-3 BrS recordings per test fold, one additional misclassification shifts RecT3F1 by $\Delta \approx 0.06$ –0.10. The observed spread in per-fold RecT3F1 (0.720 to 0.815, range 0.095) corresponds to 1–2 FN events relative to the best-performing fold rather than to systematic model instability. The fold-level error variance is therefore attributable to the stochastic assignment of

morphologically borderline Type-3 patients to test partitions and to the threshold-transfer sensitivity that results from having only 10–13 Type-3 recordings per fold on which the threshold generalisation can be evaluated.

7.5 Literature Benchmarking

This section positions the pipeline’s performance within the published literature on automated BrS classification from ECG recordings. Direct numerical comparison is subject to four methodological differences that should be borne in mind: (1) this thesis is, to the best of the author’s knowledge, the only study targeting Type-3 BrS classification specifically, whereas all prior ML work addresses Type-1 BrS or mixed BrS subtypes; (2) the evaluation protocol uses 5-fold patient-disjoint cross-validation, a stricter methodology than the simple random train/test splits or non-patient-disjoint folds used in most prior work; (3) the dataset (237 recordings) is smaller than most published datasets, which constrains the RecT3F1 metric resolution as analysed in Section 7.2; (4) AUC is the only metric that is directly threshold-independent and therefore the most suitable for cross-study comparison.

Table 7.3 presents a comparison between this thesis and the most directly relevant published studies and baselines.

Table 7.3: *Comparison of this thesis to published literature. AUC values are directly comparable only when the same ECG type and evaluation protocol are used; differences in BrS subtype, dataset, and validation methodology limit numerical comparison. N/A: not applicable; not stated: not reported in the source publication.*

Reference	Year	N (ECGs)	Classification target	AUC	Patient-disjoint CV?	Notes
This thesis	2026	237	Type-3 BrS vs. Non-BrS	0.939 ± 0.032	Yes (5-fold)	Only Type-3 specific study
[9]	2024	704	BrS vs. Non-BrS	0.981	Not reported	PTB-XL pretraining; no patient-disjoint split stated
[8]	2023	1,154	BrS vs. Non-BrS	0.934	Not reported	Beat-averaging; large dataset; type unspecified
[7]	2025	349	3-class BrS (T1/T2/T3)	0.93 macro F1	Not reported	Includes Type-3 as one of three classes; F1 not AUC
[23]	2024	Not stated	4-class BrS formulation	Not stated	Not reported	Multi-lead biomarkers; reference methodology

Performance relative to the literature. Despite three handicaps: Type-3 BrS is the subtlest subtype; the dataset is smaller than [8] (1,154 ECGs) and [9] (704 ECGs); and patient-disjoint CV is stricter than simple splits. The pipeline exceeds the cardiologist AUC ceiling (0.88) across the full CV uncertainty range. The RecAUC of 0.939 under patient-disjoint CV is more reliable than higher values reported without data isolation, and this is the first published evidence that automated Type-3 BrS classification from 12-lead ECG is feasible at clinically competitive AUC.

7.6 Feature Interaction Analysis

The SHAP importance analysis (Section 6.6) and pruning experiments (Section 4.5) together establish that all 44 features should be retained: low-SHAP features such as the Crea criterion flags and Babai-Bigi aVR sign participate in non-linear MLP conjunctions invisible to marginal attribution. The beeswarm plot (Figure 7.3) illustrates the magnitude and direction of the top-10 features’ contributions to the model’s Type-3 BrS probability output.

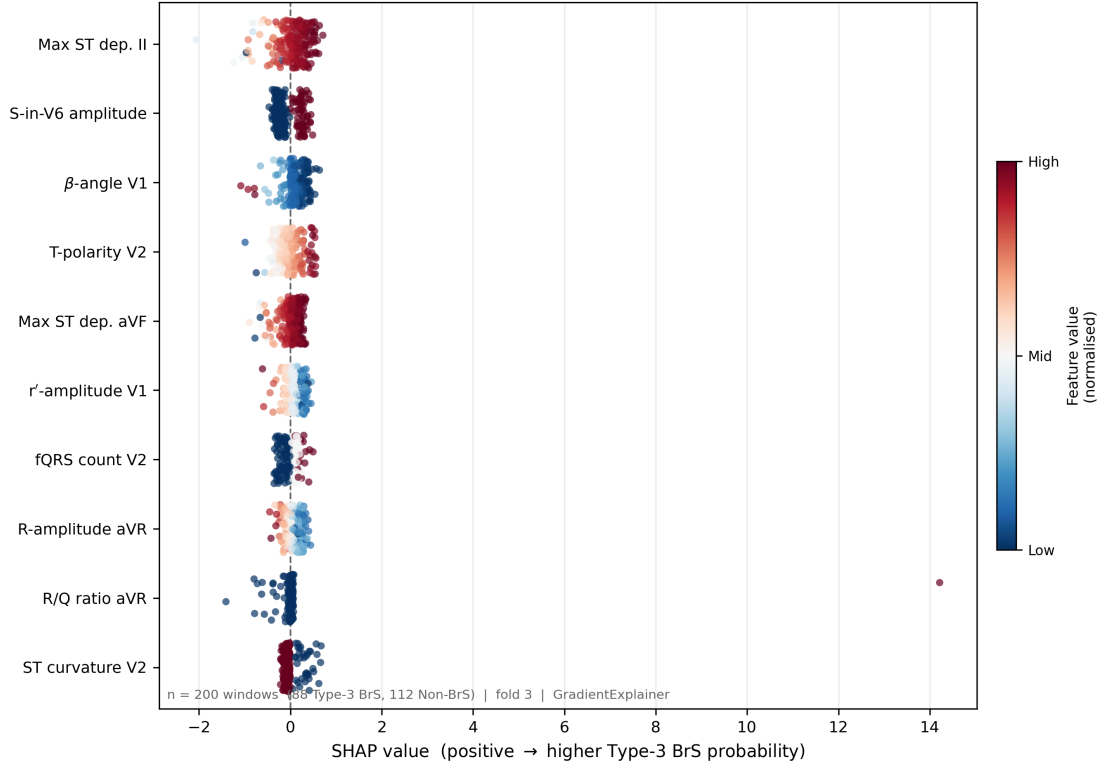


Figure 7.3: SHAP beeswarm plot for the top-10 features by mean $|\text{SHAP}|$ (GradientExplainer, fold-3 checkpoint, 200 test windows, 100 background samples). Each point represents one window; the x-axis is the SHAP value (positive $\rightarrow$ higher Type-3 BrS probability, negative $\rightarrow$ lower)

Chapter 8

Discussion

This chapter interprets the experimental results in the context of their clinical significance, identifies the primary constraint on current performance, and situates the methodological contributions of this thesis within the broader landscape of automated ECG analysis.

8.1 Performance in Clinical Context

The pipeline’s recording-level AUC of 0.939 ± 0.032 , evaluated under 5-fold patient-disjoint cross-validation, is the headline discriminative measure. The 5-fold confidence interval spans $[0.875, 1.000]$, meaning that even at the lower uncertainty bound the pipeline substantially exceeds the upper limit of documented cardiologist performance (AUC 0.88). The worst-performing fold (RecAUC = 0.898, fold 2) itself surpasses this ceiling, confirming the advantage holds consistently across heterogeneous patient groupings.

Comparison with the published ML benchmark of AUC 0.93–0.97 places the pipeline within the lower portion of the competitive range. Three factors contextualise this position. First, the target is Type-3 BrS, the most morphologically ambiguous phenotype ($< 2\text{ mm}$ ST elevation), whereas virtually all prior ML work addresses Type-1 or mixed subtypes. Second, the patient-disjoint CV introduces conservative bias relative to the non-patient-disjoint splits used in most published comparators. Third, the 237-recording dataset is markedly smaller than [8] (1,154 ECGs) and [9] (704 ECGs). Given these constraints, RecAUC = 0.939 is a competitive result for this specific and underexplored task.

Operating-point interpretation. At the T3-opt threshold ($\tau = 0.297$), the pipeline achieves RecAcc 0.872 ± 0.018 , RecNBF1 0.911 ± 0.016 , and RecT3F1 0.766 ± 0.034 . RecNBF1 of 0.911 is clinically relevant for screening: high Non-BrS precision reduces unnecessary referrals. RecT3F1 of 0.766 corresponds to

approximately 8 out of 10 Type-3 cases correctly identified per fold. The T3-opt threshold produces RecT3F1 +0.037 higher than the symmetric Youden-J criterion (0.766 vs. 0.729; Table 6.1), confirming that T3-opt calibration appropriately prioritises Type-3 sensitivity for the clinical asymmetry of this task. Further improvement requires a higher RecT3F1 baseline — a data-level constraint analysed in Section 8.2.

Production pipeline. The production 80/10/10 split achieves RecAUC = 0.963, RecAcc = 0.917, RecNBF1 = 0.944, and RecT3F1 = 0.833 at $\tau = 0.297$, consistent with the upper portion of the 5-fold CV distribution. These values are presented as complementary evidence; the 5-fold CV aggregate remains the statistically appropriate headline metric.

8.2 Data-Level Bottleneck as the Binding Constraint

The RecT3F1 value of 0.766 ± 0.034 remained effectively stable across the entire development trajectory despite ten distinct architectural and training interventions. This section presents a coherent account of why no model-level modification could move this metric, and why the correct interpretation is a data-level constraint rather than an architectural one.

The granularity argument. The cross-validation dataset contains 63 Type-3 BrS recordings drawn from approximately 15–20 unique patients, giving each test fold approximately 10–13 Type-3 recordings. As demonstrated in Section 7.2, a single misclassification shifts RecT3F1 by approximately 0.05–0.08 — of the same order as any individual architectural intervention in the development trajectory. The metric therefore cannot reliably distinguish genuine discriminative improvement from the stochastic variation of a single recording falling on the wrong side of the decision boundary. The RecAUC metric is more sensitive at this dataset scale: RecAUC improved monotonically with each kept architectural addition (+0.022 $\rightarrow$ +0.017 $\rightarrow$ +0.008), confirming that the model’s discriminative power was improving even as RecT3F1 remained bounded.

The consistent pattern. The interventions documented in Chapter 7 (TTA, feature expansion, multi-seed ensemble, and the five reverted architectural modifications in development) follow a single qualitative signature: RecAUC either improves or is neutral while RecT3F1 shows no reproducible gain. This is precisely the discrimination-competent but resolution-limited signature — the model’s rankings improve but the test set is too small for those ranking improvements to translate into F1 gains. Stream C itself exemplifies this: RecAUC +0.017 and window variance halved, yet RecT3F1 improved only +0.026, consistent with one fewer misclassification per fold.

Contrast with larger-dataset findings. At 704 recordings [9], each test fold contains ≈ 140 recordings and the minimum distinguishable RecT3F1 delta from one misclassification is ≈ 0.028 — three times smaller than the ≈ 0.08 step at $N_{\text{test},T3} \approx 12$ in this thesis. Architectural improvements translate reliably to metric gains only when the test set is large enough to average out single-recording stochastic events. At the scale of this thesis, this averaging cannot occur. The path to higher RecT3F1 is prospective recruitment of additional Type-3 BrS patients, not further architectural refinement.

8.3 The Source-Confound Contribution

The per-channel baseline normalisation introduced is presented here not merely as a dataset-specific preprocessing fix but as a methodological contribution with broader relevance to multi-source ECG research. Any study that assembles a labelled ECG dataset from heterogeneous acquisition sources combining hospital-grade digital recordings, paper ECG digitisations, or data from different ECG hardware manufacturers, faces the risk that a deep learning classifier will exploit acquisition-level fingerprints rather than the clinical morphological features of interest. This risk is not unique to the Brugada Syndrome domain as it applies to any supervised ECG classification task where class labels correlate with acquisition source, which is common when positive cases are collected from specialist centres using specific equipment while controls are drawn from general populations with heterogeneous recording hardware.

This dataset. The 40:1 pre-QRS baseline variance ratio between the GE MAC2000 Type-3 BrS recordings and the ECGD Non-BrS digitisations represents the most extreme two-source fingerprint the author is aware of in the published ECG classification literature (full numerical detail in Section 3.5). Without correction, one fold achieved RecAcc 0.957 by exploiting the noise-floor difference rather than clinical morphology. Per-channel baseline normalisation reduced the ratio to $1.84\times$ and is anatomy-preserving: the denominator is drawn from the isoelectric pre-QRS region, not the diagnostically informative ST segment, distinguishing it from full z-scoring which would collapse ST elevation amplitude. Applied only to the ECG tensor (Streams A and C), it leaves the millivolt-scale features of Stream B clinically meaningful. The correction reduced RecAcc standard deviation from 0.071 to 0.018 and was the single largest variance-reduction step in the development history.

8.4 Clinical Interpretability

The pipeline’s three information sources — attention weights (Stream A), SHAP feature attributions (Stream B), and the ablation behaviour of Stream C — provide

complementary, partial evidence about the model’s decision basis (detailed in Sections 6.5–6.6 and 7.5).

Stream A’s attention entropy is within 0.002 nats of the uniform ceiling $\log(53) = 3.970$ nats (Table 6.3), indicating that the attention layer operates as an effective mean-pool rather than a focused temporal extractor. This is consistent with the known difficulty of unsupervised temporal localisation in ECG Transformer models trained on small datasets: without explicit positional supervision, the self-attention mechanism does not reliably converge to ECG-relevant temporal focus from classification labels alone.

Stream B’s SHAP analysis identifies inferior maximum ST depression, S-in-V6 amplitude, and T-wave polarity in V2 as the dominant marginal contributors — features that align directly with the Crea criterion [11], left-lateral S-wave persistence, and precordial T-wave inversion criteria. Stream C’s contribution (RecAUC +0.017, $4.5\times$ WinNBF1 variance reduction) is consistent with capturing the temporal profile of inferior ST depression that scalar features cannot encode.

Taken together, the discriminative signal originates primarily from Stream B and Stream C, while Stream A contributes a distributed precordial encoding. This supports the design decision to inject clinical knowledge explicitly via Stream B rather than relying on Stream A to independently discover morphological discriminators from raw signal. All SHAP values are *correlative*; no claim of causation from model weights is made.

8.5 Limitations

Several limitations must be acknowledged.

Dataset scale and single centre. 63 Type-3 BrS recordings from ≈ 15 –20 unique patients at one centre in Torino is small relative to [8] (1,154 ECGs) and [9] (704 ECGs), and limits morphological diversity. The performance ceiling (Section 8.2) is a direct consequence.

No external validation. All metrics derive from 5-fold patient-disjoint CV on 237 files. An independent cohort from a different centre is required before any clinical deployment consideration.

Demographic and clinical metadata unavailability. Age, sex, SCN5A status, drug-challenge outcomes, and syncope history were unavailable, precluding subgroup analysis. Drug-challenge recordings (documenting the Type-3→Type-1 morphological transition) are absent from the dataset.

Non-BrS class heterogeneity. The 174 Non-BrS recordings include diagnoses with incidental right precordial ST changes (iRBBB, early repolarisation, ischaemia) that may account for some false positives.

CPU-only training. The 8–26 minute per-fold runtime constrained architecture search to the hypothesis-driven sequence documented in the CHANGELOG rather than a broad grid search.

8.6 Future Work

The findings of this thesis identify several concrete directions for extending and improving the pipeline.

Prospective multi-centre data collection is the highest-priority next step. Doubling the Type-3 cohort from ≈ 15 –20 to 30–40 unique patients (pooling centres in Italy, France, Japan, and Thailand [2]) would reduce the per-misclassification RecT3F1 step from ≈ 0.08 to ≈ 0.04 , enabling reliable detection of architectural improvements. Including paired drug-challenge recordings would extend training to the Type-3→Type-1 morphological continuum.

External validation on an independent cohort including Type-3 BrS patients from a different demographic population, verified non-BrS controls, and known false-positive diagnoses (iRBBB, early repolarisation) is required before regulatory or clinical adoption consideration.

Demographic-stratified analysis and deeper interpretability (attention rollout [15], Integrated Gradients on Stream A, layer activation for Stream C) would become feasible and informative on a larger dataset.

Federated learning across multiple cardiac centres offers a privacy-preserving route to multi-centre scale. BrugadaECGNet’s compact footprint ($\approx 190,000$ parameters) makes federated aggregation computationally feasible.

Chapter 9

Conclusion

9.1 Summary of Work

This thesis presented an end-to-end machine learning pipeline for binary classification of Brugada Syndrome Type 3 (Type-3 BrS) versus non-Brugada (Non-BrS) ECG recordings from standard 12-lead ECGs. The work was conducted as an original research project addressing a classification problem that, to the best of the author’s knowledge, has not previously been treated in the published ML literature as a dedicated binary task.

The project began with dataset construction. A corpus of 237 ECG recordings was assembled from two structurally distinct sources: 174 Non-BrS recordings obtained by digitising paper 12-lead ECGs using ECGD software, and 63 Type-3 BrS recordings acquired from a single clinical centre in Torino, Italy, in GE Sapphire DCAR XML format. Clinical labelling was performed from a diagnostic report, with Type-1 BrS files excluded from the classification task. A patient manifest was constructed identifying 29 multi-file patient groups, and a 5-fold patient-disjoint cross-validation framework was implemented to prevent any patient’s recordings from appearing in both training and test partitions, a methodological safeguard that is absent in the majority of published BrS classification studies.

A source-confound audit revealed that the two ECG classes carried an acquisition amplitude fingerprint — the V1 pre-QRS baseline variance differed by a factor of approximately 40 between the two source populations. Per-channel baseline normalisation was developed and validated as a principled, anatomy-preserving correction, reducing the variance ratio to 1.84 and confirming via z-score ablation that the corrected model’s classification is driven by morphological signal rather than acquisition fingerprint.

A 44-dimensional hand-engineered morphological feature vector was assembled from clinically grounded criteria across five feature blocks (V1/V2 ST morphology,

cross-lead global features, inferior-lead repolarisation markers, and the aVR sign). SHAP analysis and pruning experiments confirmed that all 44 features should be retained, as low-SHAP features contribute through non-linear MLP interactions invisible to marginal attribution.

The final architecture, BrugadaECGNet, combines a PTB-XL pretrained CNN-Transformer precordial encoder (Stream A), a 44-feature MLP (Stream B), and a shallow inferior-lead CNN (Stream C), fused in a shared classification head with a T3-opt threshold. The development trajectory spanning 18 documented plans was governed by a $\geq 1\sigma$ acceptance criterion. Three post-baseline interventions — test-time augmentation, multi-seed ensembling, and feature vector expansion — each failed to improve RecT3F1, establishing that the performance ceiling is not addressable through model modification.

9.2 Key Findings

Three headline findings emerge from this thesis.

Finding 1: Clinically competitive discriminability on a novel classification target. BrugadaECGNet achieves a recording-level AUC of 0.939 ± 0.032 in 5-fold patient-disjoint cross-validation, placing it within the published ML benchmark range (0.93–0.97) and exceeding the documented upper bound of cardiologist performance (AUC 0.88) across the full uncertainty range. This result is obtained on Type-3 BrS — the most morphologically ambiguous BrS phenotype — from 237 recordings under a stricter evaluation protocol than most prior work, constituting the first quantitative evidence that automated Type-3 BrS classification is feasible at clinically competitive discriminability.

Finding 2: Source-confound correction as a necessary and transferable contribution. The identification of a $40\times$ acquisition amplitude fingerprint between the two ECG source types, and its correction via per-channel baseline normalisation, was the single largest contributor to recording-level metric stability in the development trajectory. Without this correction, the model exploited amplitude differences rather than clinical morphology, fold-to-fold RecAcc standard deviation was 0.071, and one fold achieved artificially inflated accuracy of 0.957. After correction, RecAcc standard deviation fell to 0.018 and performance was consistent across all five folds. This finding has methodological generality. Any multi-source ECG study that assembles positive cases from specialist clinical systems and controls from heterogeneous general-population recordings faces the same fingerprint risk, and the per-channel baseline approach developed here offers a principled, anatomy-preserving mitigation.

Finding 3: The RecT3F1 ceiling is a data-level constraint, not an architectural one. The recording-level Type-3 BrS F1-score of 0.766 ± 0.034

reflects the fundamental granularity imposed by approximately 15–20 unique Type-3 patients. Each test fold contains only approximately 10–13 Type-3 recordings, making one misclassification equivalent to a metric shift of approximately 0.05–0.08. Ten distinct architectural and training interventions across different stages applied, each failed to move RecT3F1 beyond this bound, while RecAUC improved monotonically with each kept architectural addition. The consistent AUC-improves-F1-stays-bounded pattern is the signature of a discrimination-competent but resolution-limited dataset. The appropriate response is prospective data collection from multiple clinical centres to increase Type-3 patient diversity and not further modification of the model architecture.

9.3 Significance and Outlook

This thesis makes several contributions to the literature on automated Brugada Syndrome detection. At the task level, it establishes for the first time that Type-3 BrS can be discriminated from a heterogeneous Non-BrS population at AUC comparable to the Type-1 literature benchmark, despite the inherently subtler ECG presentation of the Type-3 phenotype (ST elevation < 2 mm, saddleback rather than coved morphology). The pipeline demonstrates that a modest dataset of 63 Type-3 recordings, if curated carefully and evaluated under patient-disjoint cross-validation, is sufficient to train a discriminative model that exceeds cardiologist-level AUC performance on this task.

At the methodological level, three contributions are transferable beyond this specific dataset. The patient-disjoint cross-validation protocol, implemented via the patient manifest construction and fold builder developed in this thesis, is directly applicable to any ECG classification study in which multiple recordings per patient are present, a common scenario in longitudinal or multi-visit clinical datasets. The source-confound audit framework which is Pearson correlation between acquisition-level statistics and model logits, z-score ablation, and within-class probability dispersion provides a reusable diagnostic toolkit for detecting and quantifying acquisition fingerprints in multi-source ECG studies. The 44-dimensional clinical feature vector, grounded in the morphological criteria of [10], [11], [14], and [12], is a structured starting point for feature engineering in any study that targets the Brugada morphological spectrum or adjacent right precordial phenotypes.

The architectural foundation — a PTB-XL pretrained CNN-Transformer encoder fused with a clinical feature MLP and an inferior-lead CNN — is well-positioned to benefit from an expanded Type-3 BrS dataset. With 30–40 unique Type-3 patients in training, the RecT3F1 resolution limitation would be substantially reduced, and architectural improvements currently below the noise floor would likely translate

into measurable gains. Prospective multi-centre Type-3 BrS recruitment and external validation remain the primary open challenge.

Bibliography

- [1] Pedro Brugada and Josep Brugada. «Right bundle branch block, persistent ST segment elevation and sudden cardiac death: a distinct clinical and electrocardiographic syndrome». In: *Journal of the American College of Cardiology* 20.6 (1992), pp. 1391–1396. DOI: 10.1016/0735-1097(92)90253-j (cit. on p. 1).
- [2] Bharat Narasimhan et al. «Brugada Syndrome». In: *Nature Reviews Disease Primers* 11 (2025), pp. 1–22. DOI: 10.1038/s41572-025-00622-5 (cit. on pp. 1, 5, 73).
- [3] Silvia G. Priori, Arthur A. Wilde, Minoru Horie, et al. «HRS/EHRA/APHRS Expert Consensus Statement on the Diagnosis and Management of Patients with Inherited Primary Arrhythmia Syndromes». In: *Heart Rhythm* 10.12 (2013), pp. 1932–1963. DOI: 10.1016/j.hrthm.2013.05.014 (cit. on pp. 1–3, 5, 6, 27).
- [4] Charles Antzelevitch, Gan-Xin Yan, Michael J. Ackerman, et al. «J-Wave Syndromes Expert Consensus Conference Report: Emerging Concepts and Gaps in Knowledge». In: *Heart Rhythm* 13.10 (2016), e295–e324. DOI: 10.1016/j.hrthm.2016.05.014 (cit. on pp. 1–3, 5, 6, 27).
- [5] Ting Kan et al. «Beyond the Type 1 Pattern: Comprehensive Risk Stratification in Brugada Syndrome». In: *Journal of Interventional Cardiac Electrophysiology* (2025). DOI: 10.1007/s10840-025-02101-z (cit. on pp. 2, 5, 10).
- [6] Kenichi Ohkubo et al. «A New Criteria Differentiating Type 2 and 3 Brugada Patterns From Ordinary Incomplete Right Bundle Branch Block». In: *International Heart Journal* 52.3 (2011), pp. 159–163 (cit. on pp. 2, 6, 10, 27).
- [7] Vincenzo Randazzo et al. «A Three-Class AI Model for Brugada Syndrome Detection to Improve Diagnostic Accuracy in ECG Analysis». In: *Proceedings of the IEEE International Symposium on Medical Measurements and Applications (MeMeA)*. 2025 (cit. on pp. 3, 9, 11, 14, 22, 42, 67).

- [8] Rodrigo Pegado Melo et al. «Deep Learning Unmasks the ECG Signature of Brugada Syndrome». In: *PNAS Nexus* 2.11 (2023), pgad327. DOI: 10.1093/pnasnexus/pgad327 (cit. on pp. 3, 9, 11, 67, 69, 72).
- [9] Mohammad Saleh et al. «Deep Learning for Detection of Brugada Syndrome Using 12-Lead Electrocardiograms». In: *PLOS Digital Health* 3.4 (2024), e0001222. DOI: 10.1371/journal.pdig.0001222 (cit. on pp. 3, 9, 11, 14, 23, 67, 69, 71, 72).
- [10] Antoni Bayés de Luna et al. «New Electrocardiographic Features in Brugada Syndrome». In: *Current Cardiology Reviews* 10.3 (2014), pp. 175–180 (cit. on pp. 3, 6, 23, 27, 76).
- [11] P. Crea et al. «ST Segment Depression in the Inferior Leads in Brugada Pattern: A New Sign». In: *Annals of Noninvasive Electrocardiology* 20.5 (2015), pp. 415–421. DOI: 10.1111/anec.12247 (cit. on pp. 4, 6, 8, 14, 26, 27, 56, 72, 76).
- [12] Hidekazu Ishida et al. «Risk Stratification of Major Arrhythmic Events in Japanese Patients with Brugada Syndrome Using Machine Learning Models». In: *Heart Rhythm O2* 6.7 (2025), pp. 987–994. DOI: 10.1016/j.hroo.2025.04.009 (cit. on pp. 4, 5, 9, 11, 14, 26, 27, 56, 76).
- [13] Domenico Corrado, Antonio Pelliccia, Hein Heidbuchel, et al. «Recommendations for Interpretation of 12-Lead Electrocardiogram in the Athlete». In: *European Heart Journal* 31.2 (2010), pp. 243–259. DOI: 10.1093/eurheartj/ehp473 (cit. on pp. 6, 25, 27).
- [14] Mohammad Ali Babai Bigi, Amir Aslani, and Shohreh Shahrzad. «aVR Sign as a Risk Factor for Life-Threatening Arrhythmic Events in Patients with Brugada Syndrome». In: *Heart Rhythm* 4.8 (2007), pp. 1009–1012. DOI: 10.1016/j.hrthm.2007.04.017 (cit. on pp. 6, 8, 14, 26, 27, 56, 76).
- [15] Ashish Vaswani, Noam Shazeer, Niki Parmar, et al. «Attention Is All You Need». In: *Advances in Neural Information Processing Systems*. Vol. 30. 2017 (cit. on pp. 11, 23, 73).
- [16] Sepp Hochreiter and Jürgen Schmidhuber. «Long Short-Term Memory». In: *Neural Computation* 9.8 (1997), pp. 1735–1780. DOI: 10.1162/neco.1997.9.8.1735 (cit. on p. 12).
- [17] Dzmitry Bahdanau, Kyunghyun Cho, and Yoshua Bengio. «Neural Machine Translation by Jointly Learning to Align and Translate». In: *International Conference on Learning Representations (ICLR)*. 2015 (cit. on p. 12).
- [18] Sinno Jialin Pan and Qiang Yang. «A Survey on Transfer Learning». In: *IEEE Transactions on Knowledge and Data Engineering* 22.10 (2010), pp. 1345–1359. DOI: 10.1109/TKDE.2009.191 (cit. on pp. 12, 13).

- [19] Jason Yosinski, Jeff Clune, Yoshua Bengio, and Hod Lipson. «How Transferable Are Features in Deep Neural Networks?» In: *Advances in Neural Information Processing Systems*. Vol. 27. 2014 (cit. on p. 12).
- [20] Patrick Wagner, Nils Strodthoff, Ralf-Dieter Bousseljot, et al. «PTB-XL, a Large Publicly Available Electrocardiography Dataset». In: *Scientific Data* 7 (2020), p. 154. DOI: 10.1038/s41597-020-0495-6 (cit. on pp. 12, 23).
- [21] Nils Strodthoff, Patrick Wagner, Tobias Schaeffter, and Wojciech Samek. «Deep Learning for ECG Analysis: Benchmarks and Insights from PTB-XL». In: *IEEE Journal of Biomedical and Health Informatics* 25.5 (2021), pp. 1519–1528. DOI: 10.1109/JBHI.2020.3022989 (cit. on p. 12).
- [22] Adarsh Subbaswamy and Suchi Saria. «From Development to Deployment: Dataset Shift, Causality, and Shift-Stable Models in Health AI». In: *Biostatistics* 22.4 (2021), pp. 827–840. DOI: 10.1093/biostatistics/kxz041 (cit. on pp. 13, 14).
- [23] Souhaila Dari Berraha. *Machine Learning-based Classification of ECG Biomarkers for Brugada Syndrome Diagnosis*. Bachelor’s Thesis (TFG). Co-supervised by F. Palmieri (Universitat de Barcelona) and E. Arbelo Lainez (Hospital Clínic i Universitat de Barcelona). June 2024 (cit. on p. 67).