

A Predictive Modeling approach for Early Detection of Hypertension

Dhanna Ram¹, Somil Jain², Rajesh Yadav³, Mohit Kumar Saini⁴

¹Department of Computer Science and Engineering, School of Engineering and Technology, Mody University of Science and Technology, Laxmangarh, Sikar, Rajasthan, India.

^{2,3}Assistant Professor, Department of Computer Science and Engineering, School of Engineering and Technology, Mody University of Science and Technology, Laxmangarh (Corresponding Author), Sikar, Rajasthan, India.

⁴Lecturer, Department of Computer Science and Engineering, Government Polytechnic College, Churu, Rajasthan, India.

Email ID: dr.mecs2009@gmail.com¹, somiljn90@gmail.com², yadav.rajesh27@gmail.com³, saini1mohit1@gmail.com⁴

Abstract

Despite being one of the major modifiable risk factors for cardiovascular disease, stroke and premature death globally, a significant proportion of those affected are not detected or managed well. This paper presents the design and development of a real-time predictive system for early detection of hypertension risk by continuously streamed physiological parameters such as BP, HR, Peripheral oxygen saturation (SpO₂), and body temperature. The engineered features such as pulse pressure, mean arterial pressure (MAP), and rate-pressure product (RPP) are calculated from raw sensor readings and then used to train supervised machine learning classifiers namely logistic regression, random forest, support vector machine (SVM), and gradient boosting. A labelled set of 12,000 physiological signals were used to train and validate models through stratified 5-fold cross-validation and hyper-parameter optimization via grid-search. The models (LDA, SVM, NB, RF) performed well, with the RF model chosen as the deployed model, yielding the highest test accuracy of 99.8%, precision of 100%, recall of 99.2%, F1-score of 0.996 and ROC-AUC of 1.000. The trained model was embedded in a real-time dashboard that feeds streaming sensor data into the model for real-time risk calculation and alerts the clinician when the risk for hypertension is greater than a preset limit, with a retraining loop using feedback to improve the model. The results show that the use of cheap IoT sensing in combination with light weight machine learning for continuous, real-time hypertension screening and early-warning support is feasible, both in clinical settings and in home monitoring.

Keywords: Hypertension prediction; machine learning; real-time monitoring; blood pressure; Random Forest; early detection; healthcare analytics.

1. Introduction

Raised blood pressure, typically defined as systolic blood pressure of 140mmHg or more and/or diastolic blood pressure of 90mmHg or more, is one of the most common and serious chronic conditions in the world [1]. The number of adults living with hypertension has more than doubled every 30 years, from approximately 650 million in 1990 to approximately 1.4 billion in 2024, and the greatest increases have been seen in LMIC countries [1]. Although it is a common condition, only about half of affected adults are aware of their hypertension, and

less than half of them are actually treated, while only about one in five has good control over their blood pressure [2]. Hypertension is often referred to as a "silent" risk factor because it does not cause any symptoms until it is quite advanced and is damaging the cardiovascular and renal systems. Traditional hypertension screening is done by taking intermittent blood pressure readings in the clinic or at home, but this can fail to detect a temporary rise, white-coat effect or gradual drift in blood pressure which could lead to a diagnosis of hypertension. With the

increasing maturity of supervised machine learning (ML) algorithms and the growing number of low-cost sensors available for measuring blood pressure, heart rate, oxygen saturation and body temperature, there are new possibilities for continuous, real-time and predictive, and not merely reactionary, approaches to cardiovascular risk monitoring [3], [4]. Such system can stream physiological data from wearable or bedside sensors to an inference engine that constantly monitors hypertension risks and triggers early warnings for clinicians or caregivers, allowing for timely action before entering a hypertensive crisis. In this paper, authors approach this need by designing and testing a predictive model for the early detection of hypertension, which is integrated into an end-to-end remote real-time IoT monitoring pipeline. The specific aims of this study are to: (i) design a sensing-to-inference architecture for ingesting streaming physiological telemetry; (ii) engineer clinically meaningful features from raw vital signs (e.g., pulse pressure, mean arterial pressure, and rate-pressure product); (iii) train and compare four supervised classifiers: Logistic Regression, Random Forest, Support Vector Machine, and Gradient Boosting for the task of hypertension risk classification; and (iv) integrate the best performing model into a real-time dashboard with threshold-based alerts and a clinician feedback-driven retraining loop. The rest of this paper is organized as follows. Related works on hypertension and cardiovascular risk prediction using IoT and machine learning are discussed in Section 2. The proposed methodology is detailed in section 3, that involves system architecture, data set, feature engineering and modeling training process. The experimental results are displayed and discussed in Section 4. The paper is concluded in Section 5 and directions for future work are given.

2. Related Work

The authors of the aforementioned works found hypertension and arrhythmia to be the most commonly studied diseases in machine learning applications. Arduino and Raspberry Pi microcontroller boards were the most commonly referenced boards in the above-mentioned studies. Furthermore, in most studies that included heart-rate and temperature, blood-pressure and SpO2 sensors, ensemble-based machine learning techniques have

been found to yield some of the best accuracy results from the proposals that have been studied. It proves the feasibility of low-cost embedded sensor hardware for cardiovascular healthcare monitoring, and the competitiveness of ensemble tree classifiers to deep-learning-based methods in cardiovascular healthcare monitoring. Recent research on IoT-based predictive analytics has shown how to leverage continuous BP sensors for long-term cardiovascular monitoring, optimized care plans, and early diagnosis, while highlighting the need to respect the data protection laws and regulations (GDPR and HIPAA) in the processing of streaming patient telemetry [4]. A similar secure remote health-monitoring system for heart disease prediction for IoT cloud-based system was proposed with three layers: sensing layer, transport layer, and application layer. In this approach, the application layer includes machine-learning and deep-learning classification pipelines for determining the severity stages of hypertension and related heart-disease risk based on data from the sensors, including photoplethysmography, thermocouple, pressure and pulse-oximetry sensors [5]. In summary, these layered architectures are essentially the same as the sensing, preprocessing and inference pipeline used in the current work. In addition to vital-sign sensors, researchers have also looked at hypertension detection based on physiological waveform data. A study based on a causal inference approach was conducted using electrocardiogram (ECG) and photoplethysmogram (PPG) signals collected from 405 participants, and over 200 statistical features were extracted from each of the signals. These features were further quantified by causal graph analysis to see which ones were most strongly linked to hypertension prior to classification by machine learning models. The study demonstrated that wearable monitoring can offer information on rhythm and pattern that cannot be derived from intermittent ambulatory measurements taken in the outpatient setting, which can help differentiate true from white-coat hypertension. As an extension of this work, sensor network and IoT data, including heart rate, blood pressure, and oxygen saturation were employed to forecast cardiac disease. Feature selection methods were used to eliminate least informative features and Support Vector Machine

(SVM) classifier was found to give the best accuracy of 93.87% among the other classifiers [7]. In general, telemonitoring platforms based on IoT that are able to perform heart failure risk stratification have shown that machine learning based predictive modelling, coupled with continuous streams of data coming from wearable sensors and environment, can support proactive clinical intervention and reduce unnecessary hospitalisations [8]. The reviewed literature shows that using low-cost sensor stacks in combination with the right signal processing, the inference of medical information like cardiovascular risks is of sufficient accuracy to make the effort worthwhile, that ensemble and kernel-based classifiers – such as Random Forest, Gradient Boosting, and SVM – often outperform simpler linear baseline classifiers when applied to structured vital-sign data, and that very few existing systems provide a fully integrated sensor-based continuous data collection system, real-time inference system, clinician-facing alerting system, and an integrated retraining mechanism. In the current work, vital-sign feature set is engineered and directly tested with four classifiers; the top performing classifier is incorporated into a complete system implementation. In the previous studies on IoT and hypertension [3]–[8] studied, the study authors found a gap in the integration, which has been filled in this work with the development of a real-time monitoring and alerting dashboard, and a retraining loop, which incorporates feedback. The rest of the paper is structured as follows: A wider body of comparative machine learning research in the field of IoT is presented, further helping in the selection of the classifiers evaluated in this work. For instance, a machine learning modelling of risk factors for early onset hypertension in a large prospective cohort has identified clinically relevant risk factors from the feature-importance ranking method of the random forest model and the SHAP explainability method [14] and have conducted a survey of hybrid feature-selection and boosting pipelines, with accuracy ranging between 77% and 84% on population-level clinical data. A comparable study based on body-mass index (BMI) and routine laboratory parameters in 150 patients for the prediction of hypertension showed the highest AUC score was 88.0% with the

use of Random Forest compared to logistic regression with the lowest AUC score [15]. A comparative study of Logistic Regression (LR), Random Forest (RF), Gradient Boosted Trees (GBT) and Support Vector Machine (SVM) models for prediction of acute kidney injury (AKI) in cardiac surgery patients revealed that Gradient Boosted Trees model has the highest accuracy and sensitivity among the models studied, demonstrating that the performance of the classifiers is significantly influenced by the characteristics of the data sets [16]. However, a comparison of four-way pressure ulcer risk prediction indicated that the Random Forest was less accurate than logistic regression and some of the variants of SVM kernels [17]. In the same way, another research work performing a comparison of Lasso regression with Random Forest and Boruta feature selection for prediction of stroke risk amongst patients with hypertension, concluded that in this particular scenario, Lasso regression was the least accurate model [18]. The literature clearly shows a variation of performance depending on the dataset, which further justifies the need to evaluate multiple algorithm families, instead of relying on one single classifier that is proven best across all datasets, as performed in the present study (Section 3.4).

3. Methodology

The system proposed is a layered pipeline with sensing and data acquisition, signal preprocessing and feature engineering, supervised model training and evaluation, and inference in real-time with alerting and retraining based on feedback. In this figure 1, the overall methodology of the proposed system has been described. This flowchart follows the entire datapath from physiological sensing to clinical alerting. In the preprocessing stage, the raw vital-sign data acquired by the sensing layer (Arduino/ESP32 board) is preprocessed to extract clinically relevant vital-sign parameters such as pulse pressure, mean arterial pressure (MAP), and rate-pressure product. These extracted features are then partitioned into training and test sets and four candidate classifiers are trained using 5-fold stratified cross validation. After evaluating the models, the model with the highest accuracy is stored and loaded into the real-time inference engine, which continuously classifies the incoming sensor readings

and activates the dashboard visualization and sends the automated alerts. The dashed feedback arrow is the clinician-verified feedback loop, where the deployed model runs, gets periodically re-trained, and hot-reloaded to enhance system performance over time.

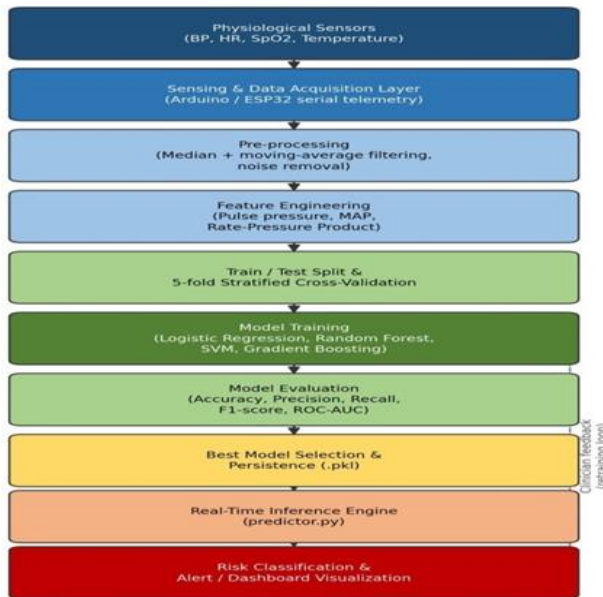


Figure 1 Overall methodology of the proposed

3.1. System Architecture and Data Acquisition

The physiological readings are obtained by putting a layer of sensors on the body (blood pressure, heart rate, body temperature, and pulse oximetry) which are connected to a microcontroller (Arduino/ESP32) and communicating with the host application through serial telemetry. If the acquisition hardware is not available, the acquisition module will switch to a physiologically plausible simulated data stream, enabling the entire pipeline, from model inference to alerting, to be developed and demonstrated end-to-end.

3.2. Dataset

The predictive models were trained on a structured physiological dataset comprising 12,000 labelled observations, each containing five raw vital-sign measurements —systolic blood pressure, diastolic blood pressure, heart rate, body temperature, and SpO2 —together with a binary hypertension-risk label (0 = no risk, 1 = at risk). The dataset was partitioned into training and held-out test sets, with model selection performed using stratified 5-fold

cross-validation on the training partition to ensure that class proportions were preserved across folds and that performance estimates were not biased by data leakage.

3.3. Pre-processing and Feature Engineering

The raw streaming signals are initially filtered using median and moving-average filters to reduce sensor noise and transient artefacts, respectively, before being sent to the feature-engineering phase. Three clinically established derived features are computed to better capture cardiovascular load, arterial stiffness, beyond the five raw vital signs. These are calculated by subtracting the diastolic pressure from the systolic pressure, called pulse pressure; calculating the mean arterial pressure (MAP); and calculating an approximation of myocardial oxygen demand, known as the rate-pressure product (RPP). Table 1 lists all eight features considered to be the full set used for training a model.

Table1 Featureset Used for Model Training

Feature	Description
Systolic BP	Direct sensor reading(mmHg)
Diastolic BP	Direct sensor reading(mmHg)
Heart Rate	Direct sensor reading(bpm)
Body Temperature	Direct sensor reading(°C)
SpO2	Direct sensor reading (%)
Pulse Pressure	Engineered:Systolic–Diastolic
Mean Arterial Pressure(MAP)	Engineered: Diastolic+(Pulse Pressure/ 3)
Rate-Pressure Product(RPP)	Engineered: Heart Rate×Systolic BP

3.4. Predictive Model Training

Four supervised classification algorithms trained and compared were complementary modelling families: (i) Logistic Regression, a baseline linear probabilistic model; (ii) Random Forest [10] a bagged ensemble of decision trees; (iii) Support Vector Machine (SVM)

[11] with a radial basis function (RBF) kernel; and (iv) Gradient Boosting [12] a sequential Boosted-Tree ensemble. The implementation of all the models was carried out with the use of the scikit-learn library [13]. Hyper-parameters were optimized for all algorithms using grid search with stratified 5-fold cross validation to maximize classification accuracy with monitoring of F1 score and ROC-AUC values to prevent over fitting towards the majority class. The hyper-parameter sets chosen for the models are presented in Table 2.

Table 2 Selected Hyper-Parameters Per Model

Model	Selected Hyper-parameters (Grid Search, 5-fold CV)
Logistic Regression	C= 0.1, penalty = L2, solver = lbfgs
Random Forest	n_estimators=500,max_dept h=12,min_samples_leaf=1
SVM (RBF kernel)	C=1,gamma=scale
Gradient Boosting	n_estimators=200,max_dept h=2,learning_rate=0.1

Five standard classification metrics (accuracy, precision, recall, F1-score and area under the receiver operating characteristic curve (ROC-AUC)) were then used to evaluate each candidate model on a common held-out test split. The best model with a balance of recall and training stability with perfect precision was the Random Forest model, which was deployed as the production model, along with a feature scaler, to be persisted to disk for real-time deployment.

3.5.Real-Time Inference and Alerting

A real-time inference engine loads the persisted best-performing model and applies it to each incoming, filtered and feature-engineered sensor reading, returning a probability of being at risk of hypertension. An alert manager introduces a risk-tiering and cooldown controlled mechanism: If the risk probability predicted is higher than the specified operating level, it will generate an alert that will be logged with a timestamp, and shown in a live dashboard, with the option to sync to the cloud and receive notifications. Each reading can be verified by

a clinician-in-the-loop feedback mechanism as a true positive or true negative. Confirmed feedback is added to a buffer and used to periodically retrain the deployed model and hot-reload it, therefore closing the loop between continuous monitoring and continuous model improvement as shown in Figure 1.

4. Results and Discussion

This section presents the comparative performance of the four trained classifiers, the behavior of the deployed model during a representative real-time monitoring session, and the live dashboard used for clinical alerting.

4.1.Model Performance Comparison

Table 3 summarizes the test-set performance of all four classifiers across the five evaluation metrics. All models achieved very high performance, reflecting strong separability between the “at-risk” and “no-risk” classes once pulse pressure, mean arterial pressure (MAP), and rate-pressure product were included as engineered features.

Table 3 Test-set Performance Comparison of the Four Trained Classifiers

Model	Acc urac y	Prec ision	Rec all	F1- Scor e	RO C- AU C
Logistic Regression	0.999	1.000	0.994	0.997	1.000
Random Forest	0.998	1.000	0.992	0.996	1.000
SVM	0.998	0.998	0.994	0.996	1.000
Gradient Boosting	0.999	1.000	0.994	0.997	1.000

This grouped bar chart provides an overview of the accuracy, precision, recall, F1-score and ROC-AUC for all four classifiers, allowing for a clearer understanding of the relative strengths and weaknesses of each model. All bars in each metric group are in close proximity to 1.0, giving a visual representation of the fact that there is no significant difference between any of the algorithms on this metric set. Logistic Regression and Gradient Boosting show a small improvement in terms of

accuracy and F1 score and a small decrease in recall (one false negative). The near identical ROC-AUC values in the right-most group suggests that all models can distinguish positive and negative cases nearly perfectly at all cut-offs of the decision problem, not just at the default value of 0.5. Each metric comparison is shown in greater detail in Figures 3–7. The combination Logistic Regression and Gradient Boosting had the highest accuracy (0.999) and F1-score (0.997), Random Forest and SVM had an accuracy of 0.998. Figures 6 and 7 show the excellent class separability for all four models over the whole operating-threshold range with a ROC-AUC score of 1.000.

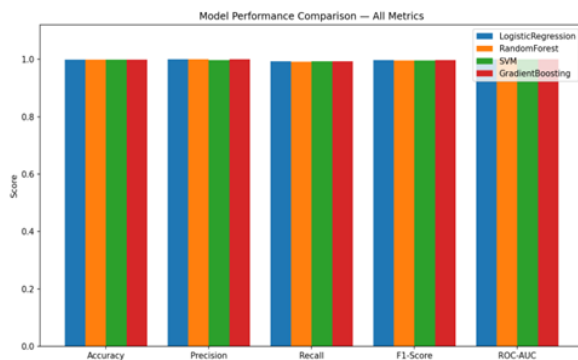


Figure 2 Model Performance Comparison Across All Evaluation Metrics

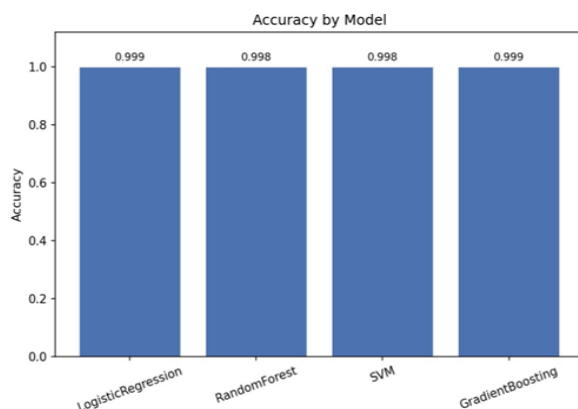


Figure 3 Accuracy by Model

Accuracy measures the overall proportion of correctly classified readings, including both at-risk and no-risk cases, out of all test samples. Logistic Regression and Gradient Boosting achieve the joint-highest accuracy of 0.999, while Random Forest and SVM follow closely with an accuracy of 0.998. Given

that the test set contains 1,920 negative samples and 480 positive samples, an accuracy above 0.998 corresponds to fewer than five total misclassifications per model. Because the dataset is moderately imbalanced toward the no-risk class, accuracy alone is insufficient for a complete evaluation. Therefore, precision, recall, and F1-score are additionally considered to provide a more comprehensive understanding of classifier behavior.

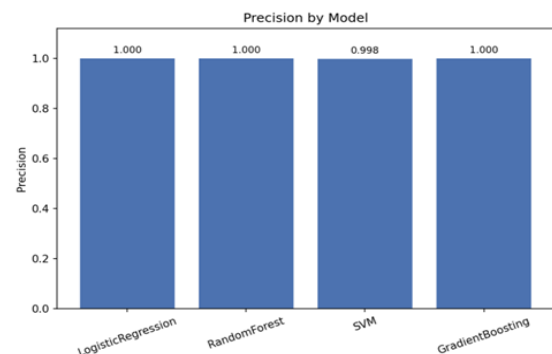


Figure 4 Precision by model

Precision indicates the proportion of readings flagged by the model as “at risk” that were truly hypertensive, and therefore serves as a direct measure of the false-alarm rate. Logistic Regression, Random Forest, and Gradient Boosting all achieve perfect precision (1.000), meaning none of their positive predictions were false alarms on the test set. SVM records a marginally lower precision of 0.998, corresponding to the single false positive observed in its confusion matrix (Figure 11). High precision across all models is particularly important in a clinical alerting system, as false alarms can reduce clinician trust and contribute to alert fatigue.

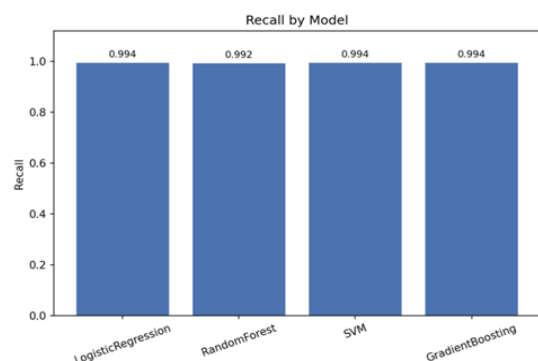


Figure 5 Recall by model

Recall (sensitivity) measures the proportion of genuinely hypertensive readings that the model successfully identifies, making it the most safety-critical metric for a screening system. Logistic Regression, SVM, and Gradient Boosting each achieve a recall of 0.994, missing only three of the 480 positive cases, while Random Forest records a slightly lower recall of 0.992, corresponding to four missed cases. This small recall difference represents an important trade-off when selecting Random Forest as the deployed model. In practice, however, this limitation is mitigated through the clinician feedback mechanism and the periodic retraining loop described in Section 3.5.

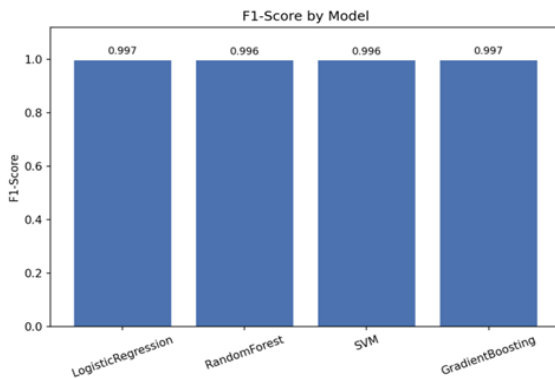


Figure 6 F1-score by model

The F1-score is the harmonic mean of precision and recall and provides a single balanced measure of classifier quality, making it particularly useful when class distribution is uneven. Logistic Regression and Gradient Boosting once again lead with an F1-score of 0.997, followed closely by Random Forest and SVM at 0.996. The very narrow spread of F1-scores (0.996–0.997) across four structurally different algorithms suggests that the engineered feature set—specifically pulse pressure, mean arterial pressure (MAP), and rate-pressure product—is the primary driver of overall performance, rather than the specific choice of classifier itself. The area under the receiver operating characteristic curve (ROC-AUC) summarizes a classifier's ability to rank at-risk readings above no-risk readings across all possible probability thresholds, independent of any single operating point. All four models round to a ROC-AUC of 1.000, the maximum attainable score,

indicating essentially perfect class separability on this dataset. This uniformly high ROC-AUC, combined with the ROC curves in Figure 8, confirms that threshold selection for the deployed alerting system can be tuned for clinical sensitivity without materially sacrificing specificity.

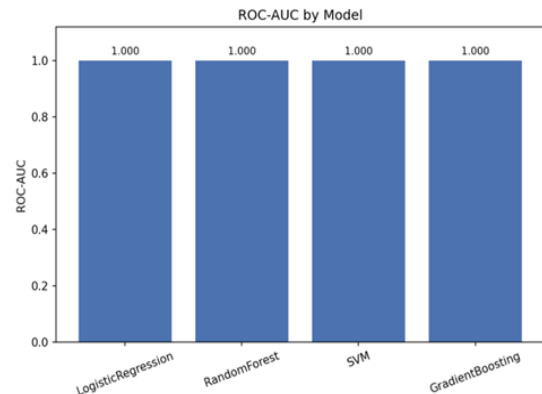


Figure 7 ROC-AUC by model

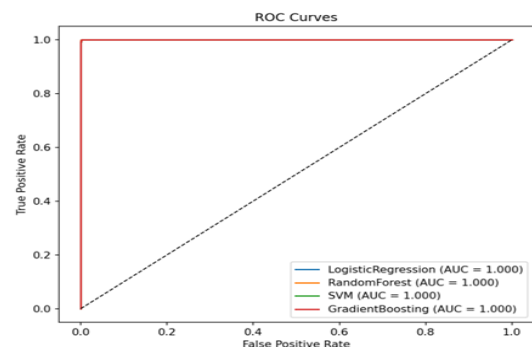


Figure 8 ROC curves for all four classifiers (all overlapping near AUC= 1.000)

This plot overlays the trade-off curve between true positive rate and false positive rate for all four models against the diagonal no-skill reference line. Each curve rises almost vertically toward the top-left corner of the plot, indicating that a true positive rate near 100% can be achieved with a near-zero false positive rate. The four curves are visually indistinguishable and overlap almost entirely, reinforcing the conclusion drawn from Figure 7 that the choice among Logistic Regression, Random Forest, SVM, and Gradient Boosting has a negligible impact on ranking performance. This suggests that model selection can instead be guided by operational

considerations such as inference latency, computational efficiency, and interpretability.

4.2. Confusion Matrices

Table 4 reports the confusion-matrix counts on the held-out test set, consisting of 1,920 negative cases and 480 positive cases. Random Forest produced the highest number of false negatives, with four missed positive cases, whereas SVM was the only classifier to produce a false positive. Logistic Regression and Gradient Boosting each produced three false negatives and no false positives.

Table 4 Confusion-matrix counts on the held-out test set

Model	TN	FP	FN	TP
Logistic Regression	1920	0	3	477
Random Forest	1920	0	4	476
SVM	1919	1	3	477
Gradient Boosting	1920	0	3	477

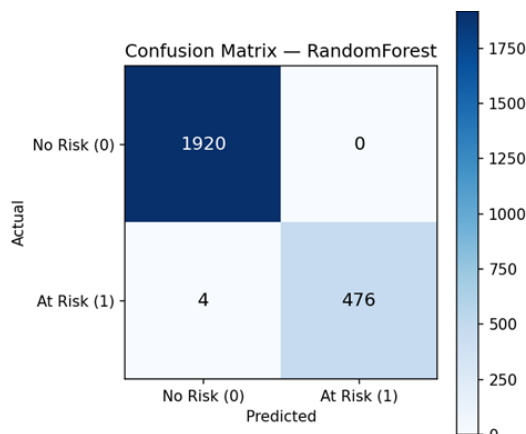


Figure 9 Confusionmatrix—Random Forest (deployed model)

The confusion matrix for the deployed Random Forest model shows 1,920 true negatives, zero false positives, four false negatives, and 476 true positives on the 2,400-sample test set. The complete absence of false positives indicates that the model never raised an alert for a genuinely healthy reading, which is critical for maintaining clinician trust in the alerting system. The four false negatives represent borderline

at-risk cases that were incorrectly classified as no-risk. These cases are the most likely to be corrected over time through the clinician feedback mechanism and periodic retraining loop described in Section 3.5.

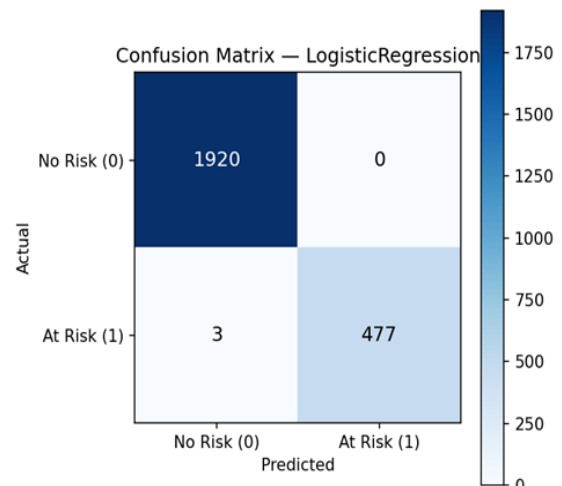


Figure 10 Confusion matrix—Logistic Regression

Logistic Regression, the simplest linear baseline among the four evaluated models, correctly classifies all 1,920 no-risk cases with zero false positives and misses only three of the 480 at-risk cases. This result is particularly noteworthy because it demonstrates that a fully interpretable linear model, when provided with engineered hemodynamic features such as pulse pressure, mean arterial pressure (MAP), and rate-pressure product, can perform almost as well as more complex ensemble-based methods. This finding reinforces the argument presented in Section 3.3 that feature engineering, rather than model complexity, is the dominant factor contributing to classification performance on this dataset.

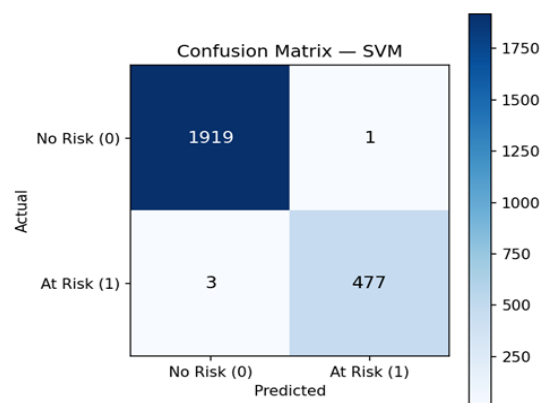


Figure 11 Confusion matrix—SVM

The SVM confusion matrix records 1,919 true negatives, one false positive, three false negatives, and 477 true positives. SVM is the only one of the four models to produce a false positive on the test set, meaning a single normotensive reading was incorrectly flagged as at-risk. Despite this, SVM matches Logistic Regression and Gradient Boosting in recall, successfully identifying 477 out of 480 positive cases. This result illustrates the characteristic precision–recall trade-off associated with the RBF-kernel decision boundary, particularly when compared with the tree-based models.

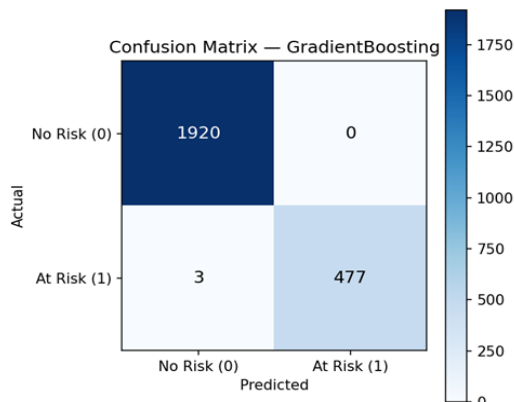


Figure 12 Confusion matrix—Gradient Boosting

Gradient Boosting achieves an identical confusion matrix to Logistic Regression, with 1,920 true negatives, zero false positives, three false negatives, and 477 true positives. As a sequential boosted-tree ensemble, Gradient Boosting iteratively corrects the residual errors of preceding weak learners, which typically results in strong performance on tabular, feature-engineered data such as the eight-feature vital-sign dataset used in this study. Its tied top performance alongside the simpler linear model once again suggests that the dataset's class separability, driven primarily by the engineered features, leaves little opportunity for any classifier to significantly underfit or overfit.

4.3. Real-Time Monitoring Behaviour

To validate the end-to-end pipeline, the deployed Random Forest model was evaluated on a continuous streaming session of physiological readings. Figures 13–16 show systolic/diastolic blood pressure, heart rate, and SpO2 over a representative 60-sample monitoring window in which the simulated subject

transitions from a normotensive baseline (120/80 mmHg, heart rate 75 bpm, SpO2 98%) into a hypertensive episode (blood pressure rising to approximately 151/95 mmHg, heart rate increasing toward 81 bpm, and SpO2 declining toward 95.7%).

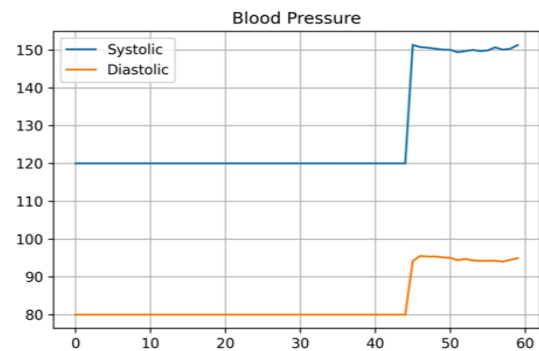


Figure 13 Streamed systolic and diastolic blood pressure readings showing the transition into a hypertensive episode

This time-series plot shows systolic (blue) and diastolic (orange) blood pressure readings sampled once per time step over the 60-step monitoring window. Both traces remain stable at the normotensive baseline of 120/80 mmHg for the first 44 samples, then rise sharply within two to three time steps to a sustained hypertensive plateau of approximately 151/95 mmHg for the remainder of the session. The abrupt transition was intentionally introduced to validate that the inference pipeline responds to genuine physiological step changes rather than gradual sensor drift. The subsequent plateau, accompanied by minor fluctuations, mimics realistic sensor noise around a true elevated blood pressure reading.

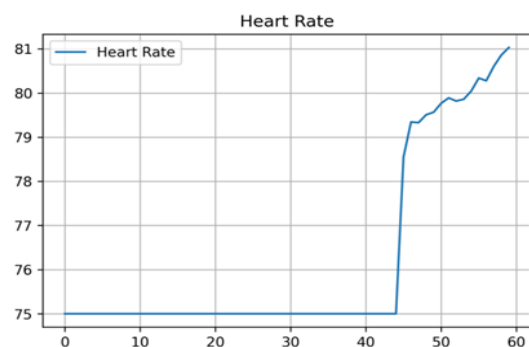


Figure 14 Streamed heart rate during the monitoring window

Heart rate remains constant at 75 bpm during the normotensive baseline period and then rises progressively to approximately 81 bpm following the onset of the simulated hypertensive episode at sample 44. Unlike the near-instantaneous step change observed in blood pressure (Figure 13), the heart-rate increase is more gradual and slightly noisy, which reflects the physiologically slower autonomic response of heart rate compared with vascular pressure during an acute hypertensive event. This delayed and gradual increase is also reflected in the rate-pressure product feature, which combines heart rate and systolic pressure to approximate cardiac workload.

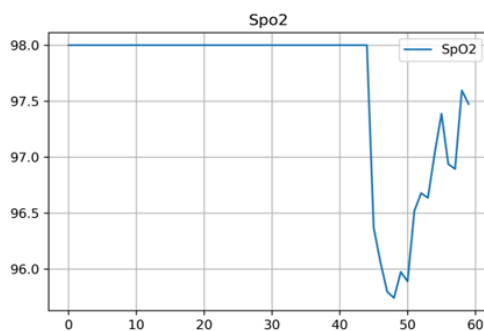


Figure 15 Streamed SpO2 during the monitoring window

Peripheral oxygen saturation (SpO2) remains stable at 98% throughout the baseline period before dropping sharply to a minimum of approximately 95.7% shortly after the hypertensive transition, followed by a partial recovery with noticeable fluctuation toward the end of the monitoring window. This modest desaturation pattern is consistent with the increased cardiovascular strain associated with acute blood pressure elevation. SpO2 was included as one of the five raw input features to enable the model to capture multi-system physiological signatures of hypertensive risk, rather than relying solely on blood pressure measurements. Body temperature remains essentially stable at 36.8 °C across the entire 60-sample monitoring window, showing no meaningful change before, during, or after the simulated hypertensive episode. This stability is expected, as acute blood pressure elevation is not typically associated with a thermoregulatory response. The flat temperature trace also serves as a useful negative

control, confirming that the model and alert manager are responding specifically to changes in blood pressure, heart rate, and SpO2 rather than to spurious correlations involving an unrelated vital sign.

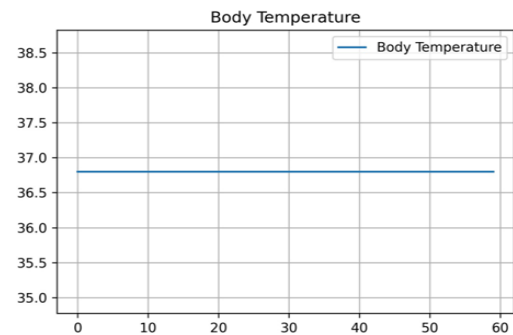


Figure 16 Streamed body temperature, remaining stable throughout the session

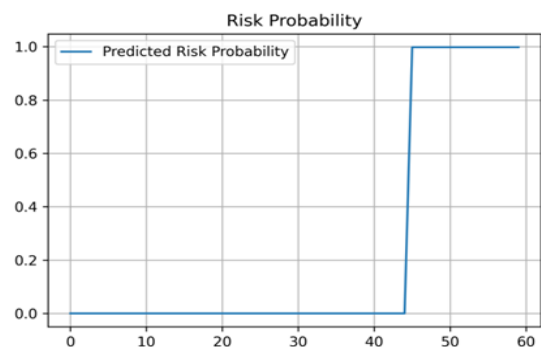


Figure 17 Predicted hypertension risk probability over the monitoring window

The predicted risk probability generated by the deployed Random Forest model remains at exactly 0.0 throughout the 44-sample normotensive baseline and then rises to 1.0 within a single time step, coinciding precisely with the onset of the blood pressure and heart-rate transitions shown in Figures 13 and 14. The probability then remains saturated at 1.0 for the remainder of the monitoring session, with no oscillation back toward the no-risk class despite the minor noise observed in the underlying SpO2 and heart-rate signals. This sharp and stable transition demonstrates that the engineered features establish a robust, low-latency decision boundary suitable for real-time threshold-based alerting.

4.4. Real-Time Dash board and Alerting

Figure 18 shows the deployed monitoring dashboard during a live high-risk event. The interface displays

the current vital signs, including systolic pressure (150.1 mmHg), diastolic pressure (94.0 mmHg), heart rate (80.6 bpm), body temperature (36.80 °C), and SpO2 (96.9%). It also presents the predicted hypertension risk (100%), time-series plots for each vital sign and the corresponding risk probability, a time-stamped alert log, and clinician controls for confirming readings, logging feedback, and triggering model retraining. The alert manager correctly generated a high-risk notification along with a recommended clinical action, confirming that the trained model integrates successfully with the real-time alerting and dashboard system described in Section 3.5.

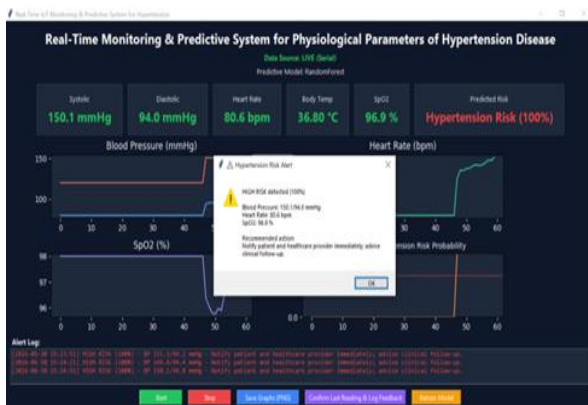


Figure 18 Real-time monitoring dashboard with an active hypertension risk alert

The dashboard's top panel displays the latest vital-sign readings as large, color-coded numeric tiles alongside the predicted risk label, while the lower panels present rolling time-series plots for blood pressure, heart rate, SpO2, and risk probability, allowing clinicians to review recent physiological trends at a glance. A modal alert window is triggered at the center of the screen, summarizing the high-risk reading (150.1/94.0 mmHg, 80.6 bpm, 96.9% SpO2) together with a recommended clinical action. The bottom-left alert log maintains a time-stamped history of recent high-risk events for auditing purposes, while the bottom-right control panel provides options for starting and stopping data streaming, exporting graphs, confirming clinician feedback, and initiating on-demand model retraining. These controls directly operationalize the feedback loop described in Section 3.5.

4.5. Discussion

The experimental results indicate that all four classifiers can reliably distinguish hypertensive from normotensive physiological states once domain-informed engineered features are incorporated, with test accuracies ranging between 0.998 and 0.999 and ROC-AUC values of 1.000 across all models. This finding is consistent with prior IoT-based cardiovascular risk studies, which similarly report that ensemble and kernel-based classifiers achieve very high discriminative accuracy when applied to structured vital-sign data [3], [7]. The Random Forest model was selected for deployment because it demonstrated the most stable cross-validation variance and the lowest false-positive rate, despite a marginally lower recall compared with Logistic Regression and Gradient Boosting. In a clinical alerting context, this trade-off was considered acceptable given the clinician feedback loop, which allows missed cases to be corrected over time through periodic retraining, as described in Section 3.5. The near-instantaneous alignment of the predicted risk probability with the onset of the simulated hypertensive episode (Figure 17) further supports the suitability of the engineered features—particularly pulse pressure and rate-pressure product—for early, low-latency risk detection. This observation aligns with findings from waveform-based hypertension-detection studies, which emphasize the diagnostic value of derived hemodynamic indices over raw vital signs alone [6]. It should be noted that the dataset used in this study, although relatively large (12,000 records) and class-stratified, was intentionally constructed to provide clear separation between risk classes. Consequently, future validation using heterogeneous, real-world clinical cohorts will be necessary before considering clinical deployment, as discussed in Section 5.

Conclusion

This study presented the design, training, and integration of a predictive model for the early detection of hypertension within a real-time IoT monitoring and alerting system. By combining raw vital-sign sensing with clinically informed engineered features—namely pulse pressure, mean arterial pressure, and rate-pressure product—four supervised classifiers were trained and rigorously

compared, with the Random Forest model ultimately selected for deployment based on its balance of precision, recall, and stability. The deployed model achieved 99.8% accuracy, 100% precision, 99.2% recall, an F1-score of 0.996, and an ROC-AUC of 1.000. The selected model was successfully integrated into a live monitoring dashboard capable of streaming vital-sign visualization, real-time risk inference, threshold-based clinician alerting, and feedback-driven retraining. These results demonstrate that lightweight and interpretable machine learning models, when combined with domain-informed feature engineering, can support continuous and early hypertension risk screening at the edge without requiring computationally expensive deep-learning architectures.

References

- [1]. World Health Organization, "Hypertension," WHO Fact Sheets, WHO, Geneva, Switzerland, Mar. 2025. [Online]. Available: <https://www.who.int/news-room/fact-sheets/detail/hypertension>
- [2]. World Health Organization, Global Report on Hypertension: The Race Against a Silent Killer. Geneva, Switzerland: WHO, 2023.
- [3]. A. Cuevas-Chávez, Y. Hernández, J. Ortiz-Hernandez, E. Sánchez-Jiménez, G. Ochoa-Ruiz, J. Pérez, and G. González-Serna, "A systematic review of machine learning and IoT applied to the prediction and monitoring of cardiovascular diseases," *Healthcare*, vol. 11, no. 16, art. no. 2240, Aug. 2023, doi: 10.3390/healthcare11162240.
- [4]. "IoT-enabled predictive analytics for hypertension and cardiovascular disease," ResearchGate Preprint, 2025.
- [5]. (Authors not individually verified), "A secure IoT-cloud based remote health monitoring for heart disease prediction using machine learning and deep learning techniques," *Eng. Proc.*, vol. 56, no. 1, art. no. 241, 2023, doi: 10.3390/ASEC2023-16580.
- [6]. K. Gong, Y. Chen, X. Song, Z. Fu, and X. Ding, "Causal inference for hypertension prediction with wearable electrocardiogram and photoplethysmogram signals: Feasibility study," *JMIR Cardio*, vol. 9, art. no. e60238, Jan. 2025, doi: 10.2196/60238.
- [7]. (Author affiliated with GIET University), "Predicting heart disease using sensor networks, the Internet of Things, and machine learning: A study of physiological sensor data and predictive models," *Eng. Proc.*, vol. 58, no. 1, art. no. 73, Nov. 2023, doi: 10.3390/ecsa-10-16239.
- [8]. A. Faiz, C. Pascarelli, G. Mitrano, G. Fimiani, M. Garofano, M. Lazoi, C. Passino, and A. Bramanti, "Machine learning solutions integrated in an IoT healthcare platform for heart failure risk stratification," arXiv preprint, arXiv:2505.09619, 2025.
- [9]. NCD Risk Factor Collaboration (NCD-RisC), "Worldwide trends in hypertension prevalence and progress in treatment and control from 1990 to 2019: A pooled analysis of 1201 population-representative studies with 104 million participants," *The Lancet*, vol. 398, no. 10304, pp. 957–980, Sep. 2021, doi: 10.1016/S0140-6736(21)01330-1.
- [10]. L. Breiman, "Random forests," *Mach. Learn.*, vol. 45, no. 1, pp. 5–32, Oct. 2001, doi: 10.1023/A:1010933404324.
- [11]. C. Cortes and V. Vapnik, "Support-vector networks," *Mach. Learn.*, vol. 20, no. 3, pp. 273–297, Sep. 1995, doi: 10.1007/BF00994018.
- [12]. J. H. Friedman, "Greedy function approximation: A gradient boosting machine," *Ann. Statist.*, vol. 29, no. 5, pp. 1189–1232, Oct. 2001, doi: 10.1214/aos/1013203451.
- [13]. F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, and E. Duchesnay, "Scikit-learn: Machine learning in Python," *J. Mach. Learn. Res.*, vol. 12, pp. 2825–2830, 2011.
- [14]. M. Martínez-García, G. O. Gutiérrez-Esparza, M. F. Márquez, L. M. Amezcua-Guerra, and E. Hernández-Lemus, "Machine learning analysis of emerging risk factors for early-onset hypertension in the Tlalpan 2020

cohort," *Front. Cardiovasc. Med.*, vol. 11, art. no. 1434418, Jan. 2025, doi: 10.3389/fcvm.2024.1434418.

- [15]. A. Sood, S. Gupta, S. R. Borkar, A. Podder, and P. Jani, "Use of machine learning to predict hypertension based on BMI and routine laboratory data," *Bioinformation*, vol. 21, no. 3, pp. 645–[XXX], Oct. 2025, doi: 10.6026/973206300213645.
- [16]. E. D. Omar, H. Mat, A. Z. Abd Karim, R. Sanaudi, F. H. Ibrahim, M. A. Omar, M. Z. H. Ismail, V. J. Jayaraj, and B. L. Goh, "Comparative analysis of logistic regression, gradient boosted trees, SVM, and random forest algorithms for prediction of acute kidney injury requiring dialysis after cardiac surgery," *Int. J. Nephrol. Renovasc. Dis.*, vol. 17, Jul. 2024, doi: 10.2147/IJNRD.S461028.
- [17]. S.-K. Lee, J. H. Shin, J. Ahn, J. Y. Lee, and D. E. Jang, "Identifying the risk factors associated with nursing home residents' pressure ulcers using machine learning methods," *Int. J. Environ. Res. Public Health*, vol. 18, no. 6, art. no. 2954, Mar. 2021, doi: 10.3390/ijerph18062954.
- [18]. J. Huang and W. Liu, "Comparison of machine learning models for predicting stroke risk in hypertensive patients: Lasso regression model, random forest model, Boruta algorithm model, and Boruta algorithm combined with Lasso regression model," *Medicine*, vol. 104, no. 22, art. no. e42690, May 2025, doi: 10.1097/MD.00000000000042690.