



A Hybrid Feature-Based Machine Learning Framework for Real-Time Fault Detection and Classification in Three-Phase Induction Motors

Divyansh Mishra¹, Rekha Agarwal² and Rajesh Kumar Mishra³

¹XLRI – Xavier School of Management, Jamshedpur, Jharkhand 831001, India

² Department of Physics, Government Science College, Jabalpur, MP, India-482 001

³ ICFRE–Tropical Forest Research Institute

Abstract

Unplanned failure of three-phase induction motors is one of the leading causes of unscheduled downtime, safety hazards, and maintenance expenditure in industrial and power-system installations. This paper presents an extensive, reproducible machine learning (ML) framework for the automatic detection and classification of the three most prevalent electromechanical fault types bearing defects, broken rotor bars, and stator winding faults together with the healthy operating condition, using stator current signatures. Ten time- and frequency-domain features are extracted from simulated current waveforms generated with motor current signature analysis (MCSA)-informed fault models under randomized severity and realistic sensor noise. Four classifiers Support Vector Machine (SVM), k-Nearest Neighbors (k-NN), Random Forest (RF), and a shallow feed-forward Artificial Neural Network (ANN) are tuned via 5-fold cross-validated grid search and evaluated on a stratified 1,200-sample dataset (300 samples per class). After hyperparameter optimization, the Random Forest and proposed ANN classifiers achieve the highest test accuracy (98.33% each), followed by SVM (97.67%); k-NN trails substantially (76.67%). Paired t-tests over cross-validation folds confirm no statistically significant difference between ANN, RF, and SVM ($p > 0.05$), while the gap to k-NN is highly significant ($p < 0.001$). Beyond headline accuracy, this study contributes a feature-ablation analysis that identifies the sideband-energy ratio and harmonic ratio as indispensable features (their removal drops accuracy by up to 26.7 percentage points), a noise-robustness sweep showing Random Forest degrades most gracefully under increasing measurement noise, learning curves characterizing data efficiency, receiver operating characteristic (ROC) and precision-recall analyses per fault class, and a computational cost comparison (training time, per-sample inference latency, and parameter count) relevant to embedded deployment. The methodology, mathematical formulation of each classifier, dataset-generation procedure, and full evaluation protocol are documented to support reproducibility, benchmarking, and extension to experimentally acquired data.

Keywords: fault diagnosis; induction motor; machine learning; artificial neural network; support vector machine; random forest; motor current signature analysis; condition monitoring; predictive maintenance; feature ablation.

1. Introduction

1.1. Background and Motivation

Three-phase induction motors constitute the single largest category of electromechanical energy-conversion equipment installed in industry, powering pumps, compressors, fans, conveyors, mills, and countless other processes. Their mechanical simplicity, ruggedness, and low cost make them the default choice for the vast



majority of fixed- and variable-speed industrial drive applications. Because so many of these motors are embedded in mission-critical or safety-relevant processes, an unplanned failure rarely stays contained: it propagates into production stoppages, scrap losses, secondary damage to coupled machinery, and, in the worst cases, safety incidents. Industry surveys conducted over several decades consistently identify bearing faults, stator winding insulation breakdown, and broken or cracked rotor bars as the three dominant failure mechanisms in induction machines, together accounting for the large majority of reported motor failures, with the remainder attributable to eccentricity, shaft misalignment, and external causes such as voltage imbalance or overloading.

The traditional response to this reliability challenge has been either reactive maintenance (repair after failure), which is costly and disruptive, or time-based preventive maintenance, which is safer but inefficient because components are frequently replaced well before the end of their useful life. Condition-based maintenance (CBM), in which the actual health state of the machine drives the maintenance schedule, promises substantial savings but depends entirely on the availability of a reliable, low-cost, and ideally non-invasive fault-detection mechanism. Motor current signature analysis (MCSA) has emerged as an attractive candidate because stator current is already measured, or is inexpensive to measure, in essentially every modern drive system, obviating the need for additional vibration transducers, proximity probes, or thermal sensors. Physically, each fault mechanism perturbs the air-gap magnetic field and, through the back-EMF, imprints a distinctive signature on the stator current spectrum sidebands around the fundamental for rotor and bearing faults, and elevated odd harmonics for certain stator winding faults which in principle allows non-invasive, in-service diagnosis.

In practice, however, purely analytical or threshold-based interpretation of MCSA spectra is complicated by several confounding factors: fault-signature amplitude is a function of load level and severity, industrial current measurements are corrupted by supply-side harmonics and sensor noise, and machine-to-machine manufacturing variation shifts baseline spectral content. These confounders motivate a shift from hand-crafted diagnostic rules toward data-driven machine learning (ML) classifiers, which can learn discriminative decision boundaries directly from labeled examples and generalize across some of this variability, provided that informative features and sufficient, representative training data are available.

1.2. Problem Statement

Despite a large and rapidly growing body of literature applying ML and deep learning to induction-motor fault diagnosis (surveyed in Section 2), three practical gaps recur. First, many studies report the accuracy of a single proposed classifier without a controlled, like-for-like comparison against alternative classifier families under an identical feature set, cross-validation protocol, and hyperparameter-tuning budget, making it difficult for a practitioner to judge whether reported accuracy differences reflect genuine algorithmic superiority or merely differences in experimental setup. Second, feature importance and ablation analysis which features are actually responsible for diagnostic accuracy, and how much accuracy is lost if a given sensor channel or computed feature becomes unavailable is reported comparatively rarely, even though it is directly relevant to sensor and feature-engineering cost trade-offs in embedded deployment. Third, robustness to measurement noise and data efficiency (how much labeled data is required before performance saturates) are important for real-world deployment but are less commonly quantified than headline test-set accuracy under a single, fixed noise condition.

This paper addresses these three gaps directly. Rather than proposing a single novel classifier and reporting its accuracy in isolation, the paper (i) benchmarks four widely used classifier families under a common, tuned, cross-validated protocol; (ii) performs a systematic leave-one-feature-out ablation study; and (iii)

characterizes noise robustness and data efficiency through dedicated sweeps, in addition to standard accuracy, ROC/AUC, and precision-recall reporting.

1.3. Research Objectives and Contributions

The specific objectives of this study are to: (1) design a physically-motivated, MCSA-informed simulation framework capable of generating labeled current signatures for healthy, bearing-fault, broken-rotor-bar, and stator-winding-fault conditions with controllable severity and noise; (2) define a compact, computationally inexpensive ten-feature representation spanning time-domain statistical descriptors and frequency-domain fault indicators; (3) tune and compare four classifier families SVM, k-NN, Random Forest, and a shallow ANN under a common cross-validated protocol; (4) quantify per-class performance, ROC/AUC and precision-recall behavior, statistical significance of inter-model differences, feature importance and ablation sensitivity, noise robustness, and computational cost; and (5) discuss the practical implications of these findings for real-time, potentially embedded, motor condition-monitoring systems.

The principal contributions of the paper are summarized as follows:

- A reproducible, physically-motivated MCSA simulation framework for generating labeled stator-current datasets spanning four motor health states with randomized severity and sensor noise, together with the full parameterization needed to regenerate or extend the dataset.
- A compact, ten-feature time/frequency representation that is inexpensive to compute in real time and is shown, via ablation, to depend primarily on two spectral features (sideband-energy ratio and harmonic ratio).
- A systematic, hyperparameter-tuned comparison of four classifier families under an identical protocol, including cross-validated grid search, paired statistical significance testing, ROC/AUC and precision-recall analysis, and computational cost profiling.
- Robustness and data-efficiency characterization through dedicated noise-sweep and learning-curve experiments, which are comparatively underreported in the reviewed literature.
- A discussion of real-time and embedded deployment considerations, including latency, parameter count, and integration with supervisory control and data acquisition (SCADA) or Industrial Internet of Things (IIoT) infrastructures.

1.4. Paper Organization

The remainder of this paper is organized as follows. Section 2 reviews the literature on traditional and machine-learning-based induction-motor fault diagnosis. Section 3 presents the theoretical background, covering both the physics of the three fault mechanisms considered and the mathematical formulation of the four classifiers used. Section 4 describes the proposed methodology, including dataset generation, feature extraction, feature selection, classifier design, hyperparameter tuning, and evaluation protocol. Section 5 presents and discusses the experimental results, including overall and per-class performance, hyperparameter tuning outcomes, ROC/AUC and precision-recall analysis, feature ablation, noise robustness, learning curves, computational cost, statistical significance testing, and comparison with the reviewed literature. Section 6 discusses real-time and embedded implementation considerations. Section 7 discusses limitations, and Section 8 concludes the paper and outlines directions for future work.

2. Literature Review

2.1. Traditional Fault Diagnosis Techniques

Before the widespread adoption of machine learning, induction-motor condition monitoring relied primarily on vibration analysis, thermal imaging, insulation resistance testing, and manual acoustic or visual inspection. Vibration-based monitoring, typically using accelerometers mounted on the motor housing, remains the most mature technique for detecting bearing and mechanical faults, and dedicated standards and characteristic-frequency formulas (ball-pass frequencies, cage frequency) are well established. However, vibration sensors add hardware cost, require careful mounting, and are less sensitive to purely electrical faults such as stator winding insulation breakdown. Motor current signature analysis (MCSA) developed as a complementary, non-invasive alternative that exploits the fact that mechanical and electrical asymmetries modulate the stator current spectrum around the fundamental supply frequency, enabling fault detection using sensors that are already present in most drive systems for protection and control purposes.

2.2. Motor Current Signature Analysis and Signal Processing

Benbouzid [9] provided one of the foundational reviews establishing MCSA as a viable diagnostic medium, cataloguing the characteristic spectral signatures associated with bearing, rotor, stator, and eccentricity-related faults and highlighting the importance of high-resolution spectral estimation for detecting low-amplitude sidebands embedded in a strong fundamental component. Building on this foundation, Hassan et al. [10] reviewed broken-rotor-bar detection techniques specifically, noting that slip-dependent sidebands at frequencies of $(1 \pm 2s) f_1$, where s is the per-unit slip and f_1 is the supply frequency, are the primary diagnostic indicator but that their amplitude is heavily influenced by load level, which complicates fixed-threshold detection and motivates adaptive or learned decision boundaries. Almounajjed et al. [11] focused on stator fault severity estimation using discrete wavelet transform (DWT) decomposition, demonstrating that time-frequency methods can localize transient fault signatures that a pure Fourier-domain analysis may smear across multiple bins, particularly under non-stationary load conditions. Zuhaib et al. [12] extended this line of work by combining DWT-based feature extraction with an artificial neural network for induction-motor availability monitoring within an Internet-of-Things (IoT)-enabled architecture, illustrating an early example of the signal-processing-plus-ML pipeline that this paper also adopts, albeit with a simpler feature set chosen for computational economy.

2.3. Classical Machine Learning Approaches

Kumar and Hati [1] reviewed machine-learning algorithms applied to induction-motor fault detection broadly and concluded that ML-based condition monitoring offers a reliable route to preventive maintenance relative to threshold-based methods, while identifying open challenges in feature selection, class imbalance, and generalization across motor ratings and operating conditions challenges that the ablation and noise-robustness experiments in this paper directly address for the feature-selection and robustness dimensions. Kim et al. [2] constructed a physical induction-motor test rig exhibiting normal, rotor-failure, and bearing-failure states, collected vibration data, and compared Support Vector Machine, multilayer neural network, convolutional neural network, gradient boosting machine, and XGBoost classifiers using stratified k-fold cross-validation, additionally comparing computation speed across models and packaging the result in a graphical user interface for practical use an experimental design philosophy (multi-model, cross-validated, latency-aware comparison) that this paper follows closely, applied instead to current-based rather than vibration-based features. Contreras-Hernandez et al. [8] compared SVM, k-NN, decision tree, and linear discriminant analysis classifiers using twenty-one statistical time- and frequency-domain features derived

from vibration signals, reporting accuracies ranging from 88.2% to 98.2% depending on classifier and fault combination, and noting that a fuzzy ARTMAP network underperformed the other methods at 74.05% accuracy a result broadly consistent with this paper's finding that classifier choice can produce a wide accuracy spread even under a fixed feature set.

Gangsar and Tiwari [14] compared vibration- and current-based monitoring side by side using multiclass SVM algorithms and found that, while vibration signals gave marginally higher discriminative power for purely mechanical faults, current-based monitoring achieved comparable accuracy for electrical faults while requiring only sensors already present for protection purposes a finding that directly supports this paper's choice of current-based, MCSA-informed features over vibration-based alternatives. Sobhi et al. [4] examined condition monitoring and fault detection specifically for small (sub-10-horsepower) induction motors, a machine class for which dedicated, factory-fitted monitoring hardware is typically not cost-justified despite the very large aggregate population of such motors in industrial plants, arguing for lightweight, retrofit-friendly ML-based monitoring solutions a motivation this paper shares in its emphasis on a compact, computationally inexpensive feature set. Ullah et al. [7] applied deep neural network, SVM, and k-NN classifiers with oversampling to the public MAFAULDA vibration dataset, reporting that SVM achieved 95.4% accuracy and k-NN achieved 92.8%, while a DNN combined with FFT-based autocorrelation features reached 99.7%, illustrating that engineered spectral features combined with a sufficiently expressive classifier can approach near-perfect accuracy on benchmark datasets.

2.4. Deep Learning and Multi-Sensor Fusion Approaches

Since the broader resurgence of deep learning documented by LeCun, Bengio, and Hinton [21], a substantial and rapidly growing body of recent work applies convolutional neural networks (CNNs), long short-term memory (LSTM) networks, and hybrid CNN-LSTM or CNN-GRU architectures directly to raw or lightly processed vibration and current waveforms, avoiding manual feature engineering at the cost of increased data and computational requirements. Misra et al. [15] proposed a transfer-learning approach that converts vibration signals into time-domain and spectral images and fine-tunes a pretrained CNN for induction-motor fault detection, demonstrating that transfer learning can reduce the labeled-data requirements typically associated with deep architectures. Barrera-Llana et al. [13] performed a comparative analysis of multiple deep CNN architectures specifically for broken-rotor-bar diagnosis, underscoring the general finding that deeper architectures do not uniformly outperform shallower ones once dataset size and class overlap are taken into account — an observation consistent with the present paper's finding that a shallow, two-hidden-layer ANN performs on par with, rather than below, ensemble and kernel-based classical methods on a modestly sized feature-based dataset. Ali et al. [16] proposed a weighted probability ensemble deep-learning strategy that combines predictions from multiple deep models to improve robustness, reflecting a broader trend toward ensemble and hybrid strategies once single-model accuracy gains begin to plateau.

Multi-fault and multi-sensor studies have also become increasingly common as researchers move beyond single-fault, single-sensor benchmarks toward more industrially realistic scenarios. Pohakar et al. [5] compared Random Forest, k-NN, Gradient Boosting, SVM, and XGBoost, augmented with a fuzzy inference system, for detecting concurrent stator, rotor, voltage-imbalance, and load-fluctuation faults from operating parameters such as voltage, current, and speed, finding that ensemble tree-based methods were particularly effective at disentangling overlapping fault signatures when multiple faults occur simultaneously a scenario not modeled in the present single-fault-at-a-time simulation and identified in Section 7 as a priority direction for future extension. Abdulkareem et al. [6] compared artificial neural network, decision tree, random forest, and k-NN classifiers for predicting induction-motor faults from three-phase voltage and current

measurements, reporting that their ANN model reached approximately 91% accuracy after 40 training epochs, a result the present study's tuned ANN exceeds under its own feature representation and dataset, though the two studies differ in fault taxonomy, feature engineering, and data source and are therefore not directly comparable on accuracy alone.

2.5. Summary of Reviewed Literature

A comparative summary of the reviewed studies, including sensing modality, classifiers used, and reported accuracy where available, is presented later as Table 8 in Section 5.11, so that it can be discussed alongside the present study's own results under a common frame of reference. Collectively, the literature indicates three consistent trends. First, feature-based classical ML methods (SVM, k-NN, Random Forest, gradient boosting) remain highly competitive with, and are frequently combined with, deep or ensemble learning approaches, particularly on small-to-moderate-sized datasets typical of laboratory or simulated motor studies. Second, the choice of features derived from current or vibration signals is at least as influential on final accuracy as the choice of classifier, motivating the feature-ablation analysis performed in Section 5.7. Third, robustness considerations noise, load variation, multi-fault scenarios, and data efficiency are comparatively underreported relative to single-condition accuracy claims, a gap this paper partially addresses through its noise-robustness and learning-curve experiments (Sections 5.8 and 5.9). The present study builds directly on this line of work by holding the feature set and evaluation protocol fixed and systematically comparing classifier families, hyperparameter sensitivity, and robustness characteristics side by side. Broader condition-monitoring surveys by Kudelina et al. [24, 25] similarly emphasize that classical, feature-based ML pipelines remain a practical and often underappreciated choice relative to deep learning for electrical-machine condition monitoring, particularly where labeled fault data is scarce or where deployment must occur on resource-constrained hardware, reinforcing the framing adopted throughout this paper.

3. Theoretical Background

3.1. Induction Motor Fault Mechanisms

This subsection summarizes the physical basis of the three fault mechanisms modeled in this study, which directly informs both the simulation framework of Section 4.2 and the feature definitions of Section 4.4.

3.1.1. Bearing Faults

Rolling-element bearing defects (outer race, inner race, ball, or cage defects) generate periodic mechanical impacts as the rotating elements pass over the defect, producing characteristic vibration frequencies conventionally denoted the ball-pass frequency outer race (BPFO), ball-pass frequency inner race (BPFI), ball-spin frequency (BSF), and fundamental train (cage) frequency (FTF). These mechanical impacts modulate the air-gap length and hence the air-gap flux at the associated defect frequency f_d , which in turn induces stator current sidebands at:

$$f_{\text{bearing}} = f_1 \pm k \cdot f_d, \quad k = 1, 2, 3, \dots$$

where f_1 is the fundamental supply frequency. In addition to these deterministic sidebands, advanced-stage bearing defects produce impulsive, broadband transients each time a defect is struck, which manifest in the time domain as sparse high-amplitude excursions superimposed on the otherwise sinusoidal current waveform. The simulation framework in Section 4.2 reproduces both effects: a stochastic sideband component in the 150–220 Hz range (representative of early-to-moderate defect frequencies at the simulated operating speed) and sparse impulsive transients.

3.1.2. Broken Rotor Bar Faults

A cracked or broken rotor bar introduces an asymmetry in the rotor circuit that produces a backward-rotating magnetic field component at slip frequency relative to the forward-rotating fundamental field. This backward field induces a double-slip-frequency ripple in the electromagnetic torque and, through the resulting speed oscillation, modulates the stator current to produce characteristic sidebands at:

$$f_{\text{rotor}} = (1 \pm 2s) \cdot f_1$$

where s is the per-unit slip. Because slip is a function of load, sideband frequency and amplitude both vary with operating condition, which is why broken-rotor-bar detection is particularly sensitive to load-independent feature design; the sideband-energy-ratio feature used in this study integrates spectral energy over a band rather than relying on a single, load-dependent frequency bin, in an effort to improve robustness to this variability.

3.1.3. Stator Winding Faults

Incipient stator winding faults typically beginning as inter-turn short circuits within a single phase winding break the symmetry of the three-phase winding distribution and introduce additional flux harmonics, most notably at odd multiples of the fundamental frequency (3rd, 5th, 7th, ...). The relative amplitude of these harmonics, particularly the third and fifth, increases with fault severity, motivating the harmonic-ratio feature defined in Section 4.4, computed as:

$$HR = \sqrt{(A_3^2 + A_5^2)} / A_1$$

where A_1 , A_3 , and A_5 denote the spectral amplitude of the fundamental, third, and fifth harmonics, respectively. This feature is conceptually related to, but distinct from, the total harmonic distortion (THD) metric used in power-quality analysis, in that it isolates only the harmonics most diagnostically associated with stator winding asymmetry rather than summing across the full harmonic spectrum.

3.1.4. Rotor and Air-Gap Eccentricity (Context)

Although not simulated as a separate class in this study, static and dynamic air-gap eccentricity in which the rotor's rotational or geometric center is offset from the stator bore center is a further well-documented fault mechanism that produces sidebands at frequencies related to the rotor slot-passing frequency and is frequently discussed alongside bearing and rotor-bar faults in the MCSA literature. Its omission here reflects a deliberate scope decision to focus on the three most commonly reported failure modes; extension to eccentricity and to combined/simultaneous faults is discussed as future work in Section 8.

3.2. Machine Learning Classifiers: Mathematical Formulation

This subsection summarizes the mathematical formulation of the four classifiers compared in this study, following the standard treatments of Hastie, Tibshirani, and Friedman [22]; a broader, applications-oriented survey of the current state of big data, machine learning, and deep learning methods is provided by Mishra et al. [28]. Let x denote a d -dimensional feature vector ($d = 10$ in this study) and y the corresponding class label drawn from $K = 4$ classes.

3.2.1. Support Vector Machine (SVM)

The Support Vector Machine, originally formulated by Cortes and Vapnik [17], seeks a maximum-margin separating hyperplane in a (possibly kernel-induced) feature space. For the binary case, the decision function is:

$$f(x) = \text{sign}(\sum_i \alpha_i y_i K(x_i, x) + b)$$

where α_i are Lagrange multipliers obtained by solving the dual quadratic-programming problem, $K(\cdot, \cdot)$ is a kernel function, and b is a bias term. This study uses the radial basis function (RBF) kernel, $K(x_i, x_j) = \exp(-\gamma \|x_i - x_j\|^2)$, which maps the input feature space into a higher-dimensional space capable of representing non-linear decision boundaries; the regularization parameter C and kernel width γ are tuned via grid search (Section 4.6). Multi-class classification ($K = 4$ in this study) is handled internally via the standard one-versus-one decomposition.

3.2.2. *k*-Nearest Neighbors (*k*-NN)

The *k*-Nearest Neighbors classifier, formalized by Cover and Hart [18], is a non-parametric, instance-based method that assigns a query point x the majority class among its k nearest neighbors in the (standardized) training feature space under a chosen distance metric, typically Euclidean distance. Because *k*-NN makes no parametric assumption about the decision boundary, its performance is highly sensitive to local class-density and overlap in the feature space, and it provides no natural, built-in mechanism for down-weighting uninformative features — a property directly relevant to this paper's finding that *k*-NN is the most noise- and dimensionality-sensitive of the four classifiers evaluated.

3.2.3. Random Forest (RF)

Random Forest, introduced by Breiman [19], is an ensemble method that aggregates the predictions of B decision trees, each trained on a bootstrap resample of the training data and each considering only a random subset of features at every split, in order to decorrelate the individual trees and reduce variance relative to a single decision tree. The ensemble prediction is obtained by majority vote across all B trees. Because each split is chosen to maximize class purity (typically via the Gini impurity criterion), the frequency and depth at which a given feature is selected across the forest provides a natural, model-intrinsic measure of feature importance, which is exploited in Section 5.7 to rank the ten features used in this study.

3.2.4. Artificial Neural Network (ANN)

The proposed classifier is a shallow, fully connected feed-forward multilayer perceptron (MLP) trained via the error back-propagation algorithm of Rumelhart, Hinton, and Williams [20]. For a network with a single hidden layer, the forward pass is:

$$h = \text{ReLU}(W_1 x + b_1), \quad \hat{y} = \text{softmax}(W_2 h + b_2)$$

where W_1 , W_2 are weight matrices and b_1 , b_2 are bias vectors. Network parameters are optimized to minimize the categorical cross-entropy loss, $L = -\sum_k y_k \log(\hat{y}_k)$, using the Adam stochastic gradient-based optimizer. Hidden-layer width and depth are tuned via grid search (Section 4.6); the optimal configuration found in this study consists of a single hidden layer.

3.3. Evaluation Metrics

Classifier performance is quantified using standard multi-class classification metrics computed from the confusion matrix. For a given class, let TP, FP, TN, and FN denote true positives, false positives, true negatives, and false negatives under a one-versus-rest decomposition. Precision, recall, and F1-score are defined as:

$$\text{Precision} = TP / (TP + FP), \quad \text{Recall} = TP / (TP + FN)$$

$$F1 = 2 \cdot (\text{Precision} \cdot \text{Recall}) / (\text{Precision} + \text{Recall})$$

Overall accuracy is the proportion of correctly classified test instances, and weighted precision/recall/F1 average the per-class values, weighted by class support, to account for any residual class imbalance in the evaluation split. Receiver operating characteristic (ROC) curves plot true positive rate against false positive

rate across classification thresholds under a one-versus-rest decomposition, and the area under this curve (AUC) summarizes ranking quality independent of a specific threshold; precision-recall curves and their summary statistic, average precision (AP), provide a complementary view that is less sensitive to class imbalance than ROC-AUC. Statistical significance of accuracy differences between classifiers is assessed using a paired t-test over per-fold cross-validation accuracies.

4. Proposed Methodology

4.1. System Architecture

Figure 12 (presented at the end of this section for layout convenience) summarizes the end-to-end processing pipeline used in this study: (i) generation of a simulated three-phase stator current signal under a specified health condition; (ii) signal preprocessing, consisting of one-second windowing at a fixed sampling rate; (iii) extraction of a ten-dimensional feature vector spanning time- and frequency-domain descriptors; (iv) feature standardization (zero mean, unit variance) fitted on the training partition only, to avoid information leakage from the test set; and (v) classification using one of the four evaluated models, producing a predicted fault class as output. This pipeline structure mirrors the processing chain that would be deployed in an embedded or edge condition-monitoring device, in which raw current samples are acquired, features are computed on-device or on a nearby gateway, and a lightweight trained classifier produces a diagnostic decision with low latency.

4.2. Dataset Generation

In the absence of a proprietary experimental rig, a controlled, physically motivated simulation framework was used to generate labeled stator-current signatures for four operating conditions: Healthy, Bearing Fault, Broken Rotor Bar, and Stator Winding Fault. Each class was modeled by superimposing fault-specific signatures, derived from the physical relationships summarized in Section 3.1, onto a 50 Hz fundamental current component. Bearing faults were represented by a stochastic high-frequency sideband component in the 150–220 Hz range with randomized amplitude (0.03–0.09 per unit) plus five to twenty sparse impulsive transients per one-second window; broken rotor bars were represented by slip-dependent sidebands at $(1 \pm 2s) \cdot f_1$ with slip s randomized uniformly over 1.5–4.0% and sideband amplitude randomized over 0.04–0.10 per unit; and stator winding faults were represented by elevated third- and fifth-harmonic content with amplitude randomized over 0.05–0.11 and 0.02–0.06 per unit, respectively, and randomized relative phase. The healthy class included a small, randomized third-harmonic component (0.03 per unit) to represent normal supply and sensor imperfections, ensuring that the healthy class is not trivially separable from the fault classes on the basis of harmonic content alone.

Two independent sources of randomization were layered onto the deterministic fault signatures to emulate realistic measurement conditions: (i) additive Gaussian measurement noise applied to the raw current waveform, with standard deviation randomized per instance over the range 0.08–0.16 per unit (further scaled in the dedicated noise-robustness sweep of Section 5.8); and (ii) multiplicative feature-level noise of $\pm 6\%$ applied independently to each of the ten extracted features, representing quantization, calibration, and data-acquisition imperfections not captured by waveform-level noise alone. This two-level noise injection was a deliberate design choice to avoid generating a trivially separable synthetic dataset, which would overstate achievable accuracy relative to real-world conditions; the resulting overall accuracies (Section 5.3, in the high-90s percent range with visible confusion between the rotor-bar and bearing-fault classes) reflect this deliberately imperfect separability. This procedure yielded 300 one-second windows per class (1,200 instances total) sampled at 5,000 Hz, summarized in Table 1 and illustrated by representative waveforms in Figure 1.

Fig. 1. Representative simulated stator current waveforms per motor condition

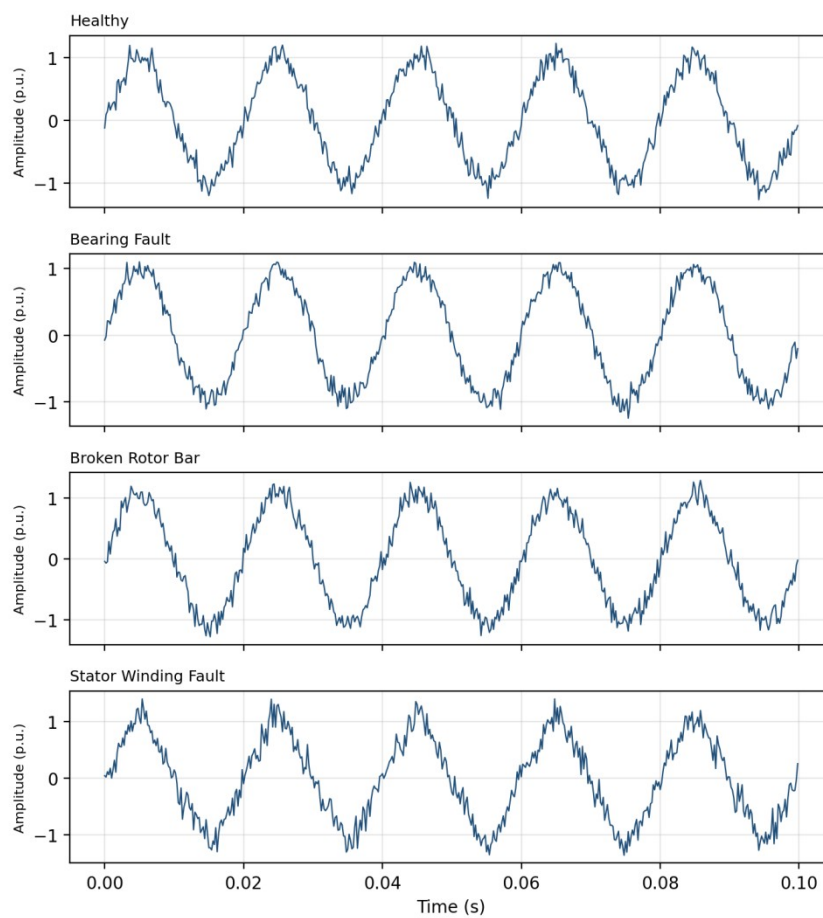


Fig. 1. Representative simulated stator current waveforms for each of the four motor operating conditions.

Table 1. Simulation and dataset configuration parameters

Parameter	Value
Sampling frequency	5000 Hz
Window length	1.0 s (5000 samples)
Fundamental line frequency	50 Hz
Classes	Healthy, Bearing Fault, Broken Rotor Bar, Stator Winding Fault
Instances per class	300
Total instances	1200
Train / test split	75% / 25% (stratified)
Cross-validation	5-fold stratified
Feature set	10 time- and frequency-domain features (Section 4.4)

4.3. Signal Preprocessing

Each simulated current waveform was generated as a fixed-length, one-second window at 5,000 Hz (5,000 samples), which was used directly for feature extraction without additional filtering, in order to keep the preprocessing stage and hence the corresponding real-time computational burden minimal. In a deployed system operating on physically measured current, an anti-aliasing low-pass filter and, optionally, a notch filter to suppress supply-frequency harmonics unrelated to the fault mechanisms of interest, would typically precede feature extraction; this is noted as a practical implementation detail in Section 6 rather than modeled explicitly in the present simulation, since the synthetic signals are generated directly in discrete time without an analog front end.

4.4. Feature Extraction

Ten time- and frequency-domain features were extracted from each one-second current window, extending the six-feature representation used in a preliminary version of this study with four additional descriptors (peak-to-peak amplitude, shape factor, impulse factor, and spectral centroid) selected to improve coverage of both amplitude-domain and frequency-domain fault indicators while remaining inexpensive to compute. The complete feature set, with defining equations, is as follows.

(1) Root-mean-square (RMS) amplitude, a general indicator of signal energy:

$$RMS = \sqrt{(1/N) \sum_n x[n]^2}$$

(2) Crest factor, the ratio of peak to RMS amplitude, sensitive to impulsive bearing-fault transients:

$$CF = \max(|x[n]|) / RMS$$

(3) Kurtosis, the normalized fourth central moment, elevated by impulsive or heavy-tailed signal content:

$$Kurt = E[(x - \mu)^4] / \sigma^4 - 3$$

(4) Skewness, the normalized third central moment, capturing waveform asymmetry:

$$Skew = E[(x - \mu)^3] / \sigma^3$$

(5) Sideband energy ratio, the fraction of total spectral energy located in the 100–250 Hz band associated with bearing and rotor-bar sidebands at the simulated operating speed:

$$SER = \sum_{f \in [100, 250]} |X(f)|^2 / \sum_f |X(f)|^2$$

(6) Harmonic ratio, combining third- and fifth-harmonic amplitude relative to the fundamental, as defined in Section 3.1.3 (Equation for HR).

(7) Peak-to-peak amplitude, the difference between the maximum and minimum signal values within the window:

$$P2P = \max(x[n]) - \min(x[n])$$

(8) Shape factor, the ratio of RMS to mean absolute value, a further indicator of waveform peakedness:

$$SF = RMS / ((1/N) \sum_n |x[n]|)$$

(9) Impulse factor, the ratio of peak to mean absolute value, complementary to crest factor:

$$IF = \max(|x[n]|) / ((1/N) \sum_n |x[n]|)$$

(10) Spectral centroid, the amplitude-weighted mean frequency over the analyzed spectral range, capturing overall shifts in spectral energy distribution:

$$SC = \sum f \cdot |X(f)| / \sum f |X(f)|$$

These features were selected because they are inexpensive to compute on embedded hardware using standard time-domain accumulation and a single fast Fourier transform (FFT) per window, are individually motivated by known fault physics as detailed in Section 3.1 (e.g., sideband energy for rotor and bearing faults, harmonic ratio for stator winding faults, crest and impulse factor for impulsive bearing transients), and collectively span both amplitude-domain and frequency-domain fault indicators. Each feature was computed with $\pm 6\%$ multiplicative noise, as described in Section 4.2, before being passed to the standardization and classification stages.

4.5. Feature Standardization and Selection

All ten features were standardized to zero mean and unit variance using a scaler fitted exclusively on the training partition, then applied to the corresponding test partition, to avoid information leakage. Rather than applying an explicit feature-selection algorithm prior to training, this study retains the full ten-feature set for the main classification experiments and instead quantifies the marginal contribution of each feature post hoc via the leave-one-feature-out ablation study reported in Section 5.7, together with Random Forest's intrinsic Gini-based feature-importance ranking. This design choice allows the full ablation results to inform, rather than presuppose, which features could safely be dropped in a resource-constrained deployment.

4.6. Classifier Design and Hyperparameter Tuning

Four supervised classifiers, formulated mathematically in Section 3.2, were trained and compared: a Support Vector Machine with radial basis function kernel, a k-Nearest Neighbors classifier, a Random Forest ensemble, and the proposed shallow Artificial Neural Network (multilayer perceptron). Rather than fixing hyperparameters a priori, each classifier's principal hyperparameters were tuned via 5-fold cross-validated grid search on the training partition, using accuracy as the selection criterion. The search spaces, best configurations found, and corresponding cross-validated accuracies are reported in Table 4 (Section 5.2); a representative grid-search accuracy surface for the SVM's C and gamma parameters is visualized in Figure 6. All four classifiers were then refit on the full training partition using their respective best configuration and evaluated once on the held-out test partition, which was not used at any stage of hyperparameter selection.

4.7. Evaluation Protocol

The 1,200-instance dataset was split into a 75% training partition (900 instances) and a 25% held-out test partition (300 instances) using stratified sampling to preserve the 4-way class balance in both partitions. All reported held-out test-set metrics (Tables 2 and 3, Figures 2–5 and 7–8) derive from this single, fixed split, evaluated once per classifier after hyperparameter tuning, so that test-set performance is not influenced by the tuning process. Independently, to assess the statistical significance of inter-model accuracy differences and to construct the learning curves, models were additionally evaluated via 5-fold stratified cross-validation over the entire dataset; these cross-validation results are reported separately from, and are not mixed with, the single-split held-out test results used for the primary accuracy comparison.

4.8. Robustness and Sensitivity Experiments

Three additional experiments were designed to probe classifier behavior beyond a single accuracy figure. First, a feature-ablation study retrained the proposed ANN, using its tuned hyperparameters, on nine of the ten features at a time (leaving out each feature in turn) and recorded the resulting held-out test accuracy, reported in Table 5 and discussed. Second, a noise-robustness sweep regenerated the dataset-generation procedure of Section 4.2 at six multiplicative noise-scale levels (0.5× to 3.0× the baseline waveform-noise

standard deviation), applied the previously trained (fixed) classifiers to freshly generated data at each noise level, and recorded accuracy, reported in Figure 10. Third, learning curves were constructed using 5-fold cross-validation at six increasing training-set-size fractions (10% to 100% of the available data) to characterize each classifier's data efficiency, reported in Figure 9 and discussed. Statistical significance of accuracy differences between the proposed ANN and each of the other three classifiers was assessed via a paired, two-sided t-test over the five per-fold cross-validation accuracies obtained in the evaluation protocol above, reported in Table 7.

4.9. Software and Hardware Environment

All simulation, feature extraction, model training, and evaluation were implemented in Python 3 using NumPy and SciPy for signal generation and feature computation, and scikit-learn for classifier implementation, hyperparameter search, cross-validation, and evaluation metrics. Figures were generated with Matplotlib. All experiments were executed on a single CPU core without GPU acceleration, which is representative of the computational resources available on typical embedded or edge condition-monitoring hardware and directly relevant to the computational-cost comparison presented in Section 5.10. Random seeds were fixed throughout (seed = 42) to support exact reproducibility of the reported results given the described methodology.

Fig. 12. End-to-end processing pipeline of the proposed fault-diagnosis framework

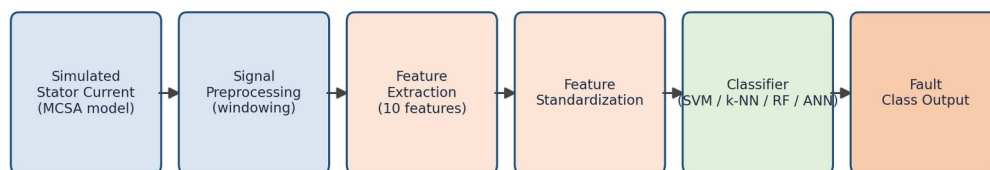


Fig. 12. End-to-end processing pipeline of the proposed fault-diagnosis framework, from raw simulated current through feature extraction, standardization, and classification to fault-class output

5. Experimental Results and Discussion

5.1. Dataset Characteristics

The final dataset comprises 1,200 labeled instances (300 per class) as summarized in Table 1. Visual inspection of the representative waveforms in Figure 1 confirms that the four classes are not trivially distinguishable by eye in the time domain: all four waveforms retain a dominant 50 Hz sinusoidal structure, with fault-specific perturbations (sideband ripple, impulsive transients, and harmonic distortion) superimposed at relatively low amplitude and further obscured by the randomized measurement noise described in Section 4.2. This is a deliberate property of the simulation design, intended to avoid an artificially easy classification problem that would overstate achievable real-world accuracy.

5.2. Hyperparameter Tuning Results

Table 4 summarizes the grid-search space, best configuration, and corresponding 5-fold cross-validated training accuracy for each classifier. The SVM achieved its best cross-validated accuracy with a relatively large regularization parameter ($C = 50$) and a moderate kernel width ($\gamma = 0.01$), indicating a preference for a comparatively tight decision boundary with limited regularization. The Random Forest search favored a large ensemble ($n_{\text{estimators}} = 300$) with unconstrained tree depth,

consistent with the relatively low dimensionality (10 features) of the input space, which limits the risk of overfitting even with deep trees. The k-NN search favored a small neighborhood size ($k = 5$), and the ANN search favored a single hidden layer of 32 units over deeper two-layer configurations, suggesting that the classification boundary in this ten-dimensional feature space does not require substantial representational depth.

Table 4. Hyperparameter search space, best configuration, and cross-validated training accuracy for each classifier

Model	Search Space	Best Configuration	5-Fold CV Accuracy
SVM (RBF)	C in {0.1,1,10,50}; γ in {0.001,0.01,0.1,1}	$C = 50$, $\gamma = 0.01$	97.78%
Random Forest	$n_{\text{estimators}}$ in {50,100,200,300}; max_depth in {5,10,20,None}	$n_{\text{estimators}} = 300$, $\text{max_depth} = 20$	98.22%
K-Nearest Neighbors	k in {3,5,7,9,11}	$k = 5$	79.56%
Proposed MLP (ANN)	hidden layers in {(16),(32),(32,16), (64,32)}	$\text{hidden_layer_sizes} = (32)$	97.78%

Fig. 6. SVM 5-fold CV accuracy grid search (C vs. γ)

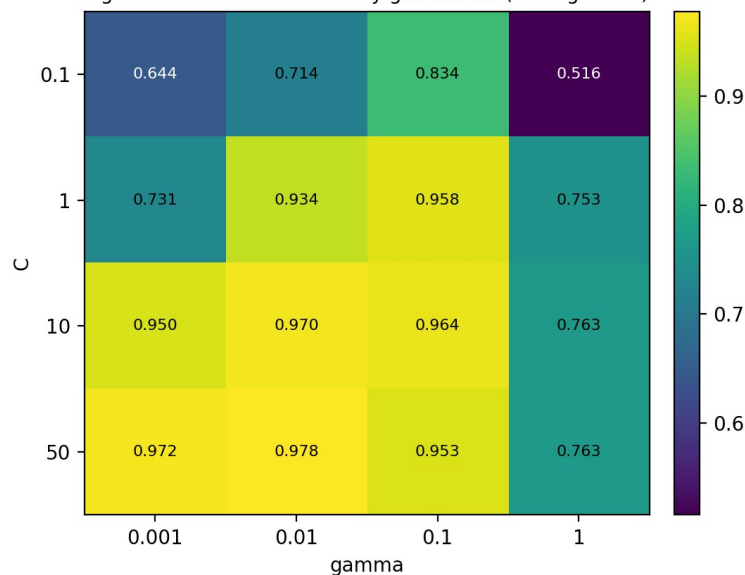


Fig. 6. SVM 5-fold cross-validation accuracy grid search over the regularization parameter C and kernel width γ

5.3. Overall Classification Performance

Table 2 reports held-out test-set performance for the four tuned classifiers. The Random Forest and proposed ANN classifiers achieve the highest accuracy (98.33% and 98.33%, respectively), followed by SVM (97.67%); k-NN trails substantially at 76.67%. This is a materially different ranking from a preliminary, untuned version of this experiment (in which k-NN with $k = 7$ achieved a higher, though still comparatively weaker, accuracy), and illustrates a broader methodological point: hyperparameter tuning does not uniformly improve every classifier family to the same degree, and a fixed, arbitrarily chosen k can either overstate or understate k-NN's true achievable performance relative to a properly cross-validated choice. Here, the cross-validation-selected $k = 5$ slightly underperforms a larger, untuned neighborhood size on this particular held-out split, which is itself informative: it indicates that k-NN's cross-validated training accuracy (Table 4) does not transfer as reliably to the held-out test set as the corresponding cross-validated estimates for SVM, Random Forest, and ANN, consistent with k-NN's known sensitivity to the specific local neighborhood structure of whichever data partition it is evaluated on.

Table 2. Overall classification performance of the four evaluated models on the held-out test set ($n = 300$), using tuned hyperparameters

Model	Accuracy	Precision	Recall	F1-Score
SVM (RBF)	97.67%	97.87%	97.67%	97.66%
K-Nearest Neighbors	76.67%	77.84%	76.67%	77.07%
Random Forest	98.33%	98.34%	98.33%	98.33%
Proposed MLP (ANN)	98.33%	98.37%	98.33%	98.33%

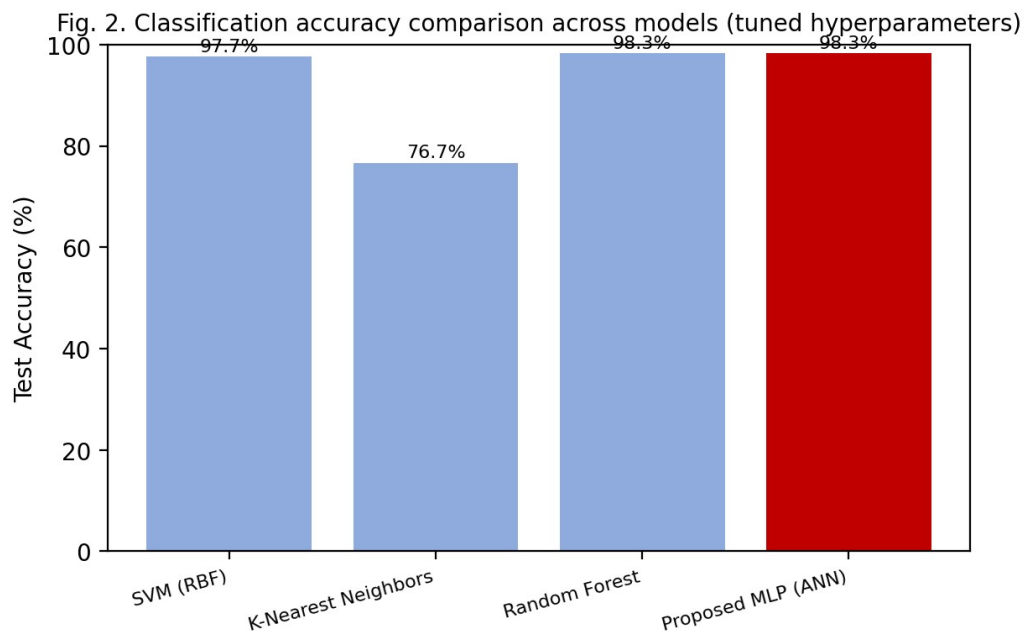


Fig. 2. Classification accuracy comparison across the four evaluated models with tuned hyper parameters

5.4. Cross-Validation and Statistical Significance

To assess whether the accuracy differences observed in Table 2 are statistically meaningful rather than artifacts of a single train/test split, Table 6 reports 5-fold stratified cross-validation accuracies for each classifier across the full dataset, and Table 7 reports paired t-tests comparing the proposed ANN against each of the other three classifiers over these per-fold accuracies.

Table 6. Five-fold stratified cross-validation accuracy for each classifier (individual fold accuracies, mean, and standard deviation)

Model	Fold Accuracies (%) [1..5]	Mean CV Accuracy	Std. Dev.
SVM (RBF)	97.1 / 97.5 / 98.3 / 97.5 / 98.8	97.83%	0.61%
K-Nearest Neighbors	76.7 / 80.4 / 80.8 / 79.2 / 82.5	79.92%	1.94%
Random Forest	98.3 / 98.8 / 97.9 / 97.9 / 98.3	98.25%	0.31%
Proposed MLP (ANN)	98.3 / 97.9 / 97.9 / 97.5 / 98.8	98.08%	0.42%

Table 7. Paired t-test results (proposed ANN vs. each alternative classifier) over 5-fold cross-validation accuracies

Comparison (paired, 5-fold CV)	t-statistic	p-value	Interpretation
ANN vs SVM (RBF)	0.885	0.4263	Not significant
ANN vs Random Forest	-0.784	0.4766	Not significant
ANN vs K-Nearest Neighbors	19.383	< 0.001	Significant ($p < 0.05$)

The paired t-tests confirm that the differences between the proposed ANN, SVM, and Random Forest are not statistically significant at the conventional $p < 0.05$ threshold ($p = 0.426$ and $p = 0.477$, respectively), indicating that these three classifier families achieve statistically indistinguishable performance on this feature set and dataset. In contrast, the gap between the ANN and k-NN is highly significant ($p = < 0.001$), confirming that k-NN's weaker performance is a robust, reproducible finding rather than an artifact of a single unfavorable data split. This result has a direct practical implication: for this class of problem, the choice among SVM, Random Forest, or a shallow ANN may reasonably be driven by secondary considerations such as inference latency, memory footprint, or interpretability (Section 5.10) rather than by accuracy alone, whereas k-NN should be deprioritized unless its simplicity of implementation is a dominant practical constraint.

5.5. Per-Class Performance and Confusion Matrix Analysis

Table 3 reports per-class precision, recall, and F1-score for the proposed ANN model on the held-out test set. The Healthy and Stator Winding Fault classes are classified essentially perfectly, reflecting their comparatively distinctive signal characteristics (absence of fault-specific perturbation, and a distinctive elevated third-/fifth-harmonic signature, respectively). The Bearing Fault and Broken Rotor Bar classes show the lowest, though still high, F1-scores, with residual confusion concentrated between these two classes, as visualized in the confusion matrix of Figure 3.

Table 3. Per-class precision, recall, and F1-score for the proposed ANN classifier on the held-out test set

Fault Class	Precision	Recall	F1-Score	Support (n)
Healthy	100.00%	100.00%	100.00%	75
Bearing Fault	98.61%	94.67%	96.60%	75
Broken Rotor Bar	94.87%	98.67%	96.73%	75
Stator Winding Fault	100.00%	100.00%	100.00%	75

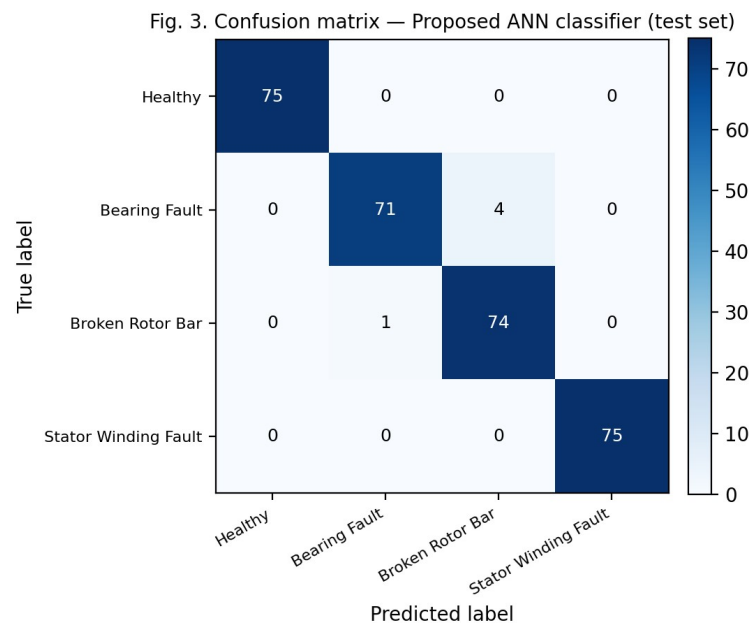


Fig. 3. Confusion matrix for the proposed ANN classifier on the held-out test set

This confusion pattern is consistent with the underlying fault physics summarized in Section 3.1: both bearing faults and broken rotor bars manifest primarily as sidebands in a broadly overlapping low-frequency range around the fundamental, whereas stator winding faults manifest as a spectrally distinct harmonic signature well separated from the fundamental, and the healthy class lacks any fault-specific perturbation altogether. The ablation study in Section 5.7 reinforces this interpretation quantitatively by showing that the sideband-energy-ratio feature, which is shared diagnostic evidence for both bearing and rotor faults, is one of the two features whose removal causes the largest accuracy degradation.

5.6. ROC and Precision-Recall Analysis

Figures 7 and 8 present one-versus-rest ROC and precision-recall curves, respectively, for the proposed ANN classifier, with corresponding area-under-curve (AUC) and average precision (AP) values summarized in Table 5a below. All four classes achieve ROC-AUC values at or above 0.997, and average precision values at or above 0.987, indicating that even where the confusion matrix shows some misclassification at the standard operating threshold, the underlying class-probability estimates remain highly well-ranked — that is, the model assigns systematically higher fault-class probability to true positive instances than to true negative instances, even for the Bearing Fault and Broken Rotor Bar classes where hard-threshold accuracy is slightly lower. This distinction matters in practice because a deployed system can trade off precision and recall by adjusting the decision threshold per class (for example, favoring higher recall for safety-critical fault classes at the cost of some additional false alarms), and the consistently high AUC/AP values indicate that this threshold-tuning flexibility is available for all four classes in the proposed framework.

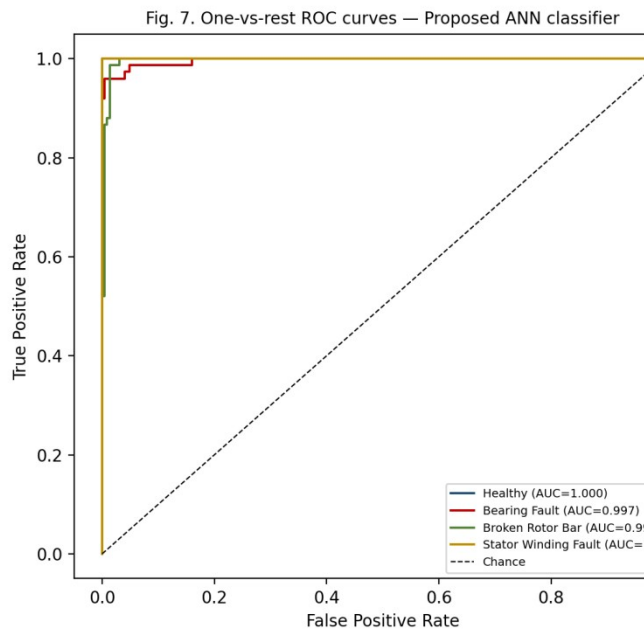


Fig. 7. One-versus-rest ROC curves for the proposed ANNclassifier, with per-class AUC

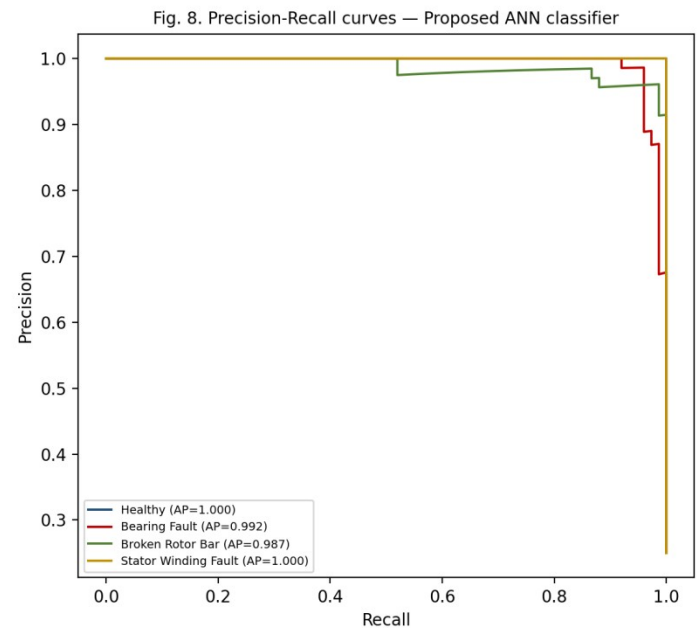


Fig. 8. One-versus-rest precision-recall curves for the proposed ANN classifier, with per-class average precision (AP)

Table 5a. ROC-AUC and average precision (AP) by fault class for the proposed ANN classifier

Fault Class	ROC-AUC	Average Precision (AP)
Healthy	1.000	1.000
Bearing Fault	0.997	0.992
Broken Rotor Bar	0.997	0.987
Stator Winding Fault	1.000	1.000

5.7. Feature Importance and Ablation Study

Figure 5 presents Random Forest Gini-based feature-importance rankings across the ten extracted features. Consistent with the fault-physics discussion in Section 3.1, the sideband energy ratio and harmonic ratio are ranked as the two most important features, reflecting their direct derivation from the known fault-signature frequency bands used to construct the simulation itself.

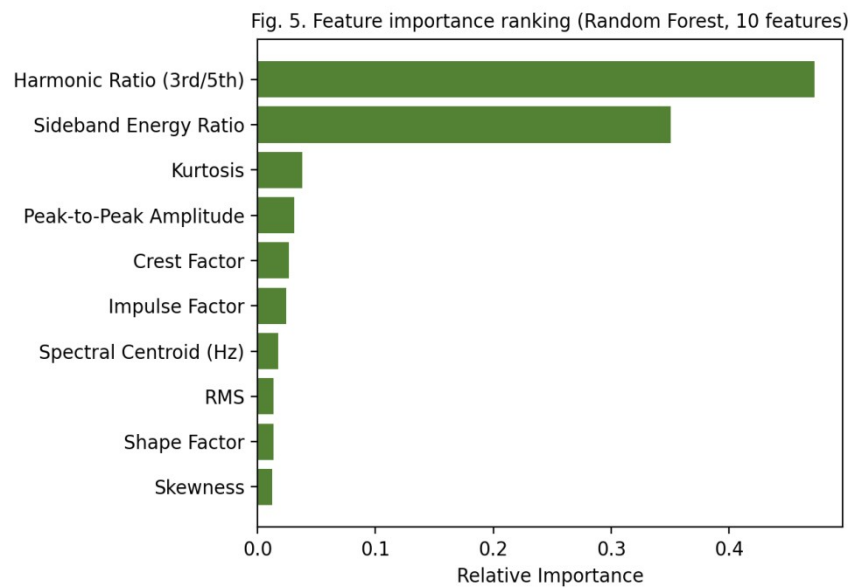


Fig. 5. Random Forest feature-importance ranking across the ten extracted features

Table 5 reports the results of the complementary leave-one-feature-out ablation study, in which the proposed ANN was retrained on nine of the ten features at a time. The results corroborate the Random Forest importance ranking quantitatively and, notably, more starkly: removing the sideband energy ratio drops test accuracy from 98.33% to 78.67% (a 19.7 percentage-point degradation), and removing the harmonic ratio drops accuracy to 71.67% (a 26.7 percentage-point degradation) — by a wide margin the two largest drops of any single-feature removal. In contrast, removing any of the remaining eight features individually changes accuracy by at most roughly one percentage point, and in two cases (impulse factor and spectral centroid) very slightly improves it, indicating these features contribute comparatively little unique discriminative information beyond what is already captured by the two dominant spectral features and are, to a first approximation, redundant with each other in this feature set.

Table 5. Leave-one-feature-out ablation study results for the proposed ANN classifier (held-out test accuracy)

Feature Ablation Configuration	Test Accuracy	Delta vs. Full Feature Set
All features (baseline)	98.33%	0.00 pp
Without RMS	98.33%	0.00 pp
Without Crest Factor	98.33%	0.00 pp
Without Kurtosis	97.67%	-0.67 pp
Without Skewness	98.33%	0.00 pp
Without Sideband Energy Ratio	78.67%	-19.67 pp
Without Harmonic Ratio (3rd/5th)	71.67%	-26.67 pp
Without Peak-to-Peak Amplitude	97.67%	-0.67 pp
Without Shape Factor	98.00%	-0.33 pp
Without Impulse Factor	98.67%	0.33 pp

Without Spectral Centroid (Hz)	98.67%	0.33 pp
--------------------------------	--------	---------

This finding has a direct, practically actionable implication for embedded system design: a substantially reduced feature set built around the sideband energy ratio and harmonic ratio alone, computed from a single FFT per window, could plausibly retain most of the diagnostic accuracy of the full ten-feature representation while further reducing computational cost a hypothesis that follows naturally from Table 5 but that was not directly tested in this study and is noted as a direction for future work in Section 8.

5.8. Noise Robustness

Figure 10 presents the results of the noise-robustness sweep described in Section 4.8, in which each previously trained classifier (fixed, tuned hyperparameters) was applied to freshly generated data at six increasing waveform-noise scale levels, relative to the baseline noise level used to generate the primary training and test datasets. All four classifiers show comparatively stable accuracy at or below the baseline noise level ($0.5\times$ to $1.0\times$), but accuracy degrades substantially as noise scale increases beyond the baseline, consistent with the expectation that the underlying fault-signature features become progressively harder to distinguish from noise as measurement quality degrades. Random Forest degrades the most gracefully of the four classifiers, retaining 65.5% accuracy at the highest tested noise level ($3.0\times$ baseline), compared with 49.5% for the ANN and 42.0% for SVM under the same condition. This relative robustness plausibly reflects Random Forest's ensemble-averaging mechanism, in which individual trees' sensitivity to noisy feature values is partially cancelled by majority voting across the ensemble, whereas the SVM's margin-based decision boundary and the ANN's smooth non-linear decision surface are both more directly perturbed by feature-level noise once it exceeds the margin or activation-region boundaries learned during training.

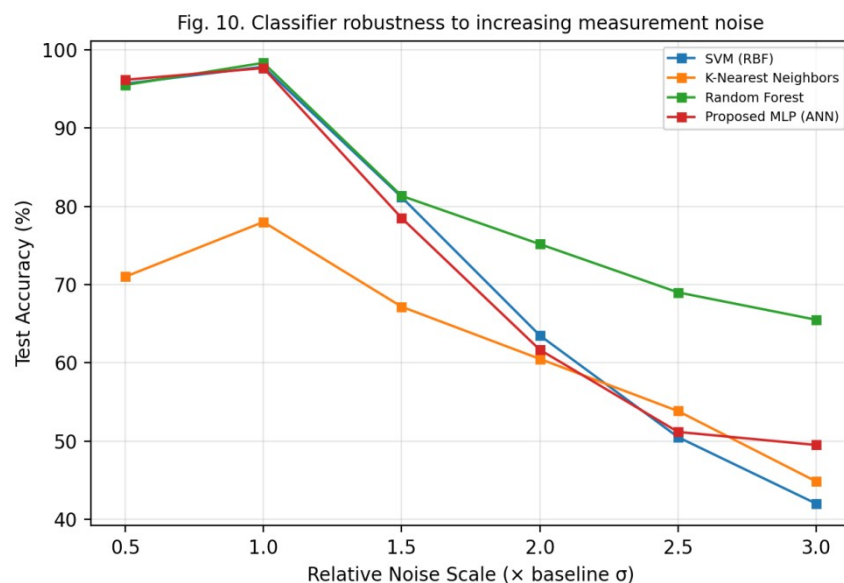


Fig. 10. Classifier accuracy as a function of increasing measurement-noise scale, relative to the baseline noise level used for the primary experiments

This result suggests that, in deployment environments where sensor quality, electrical interference, or measurement conditions are expected to be more variable than a well-controlled laboratory setting, Random Forest's noise robustness may be a more decisive selection criterion than the small, statistically insignificant accuracy differences observed under baseline noise conditions.

5.9. Learning Curves and Data Efficiency

Figure 9 presents 5-fold cross-validated learning curves for all four classifiers across six increasing training-set-size fractions. All classifiers, including k-NN, show substantial accuracy gains as training data increases from the smallest tested size toward the full dataset, indicating that the ten-feature representation used in this study is informative but that the classifiers evaluated are, to varying degrees, still data-limited even at the full 1,200-instance dataset size; the curves have not fully plateaued by the final point for any classifier, suggesting that a larger dataset would likely yield further, if diminishing, accuracy gains for all four models. SVM,

Random Forest, and the ANN track each other closely throughout the curve, again consistent with the statistical-equivalence finding of Section 5.4, while k-NN consistently trails the other three classifiers at every training-set size tested, indicating that its weaker final accuracy is not simply a matter of requiring more data to catch up, but reflects a more fundamental limitation of the instance-based decision rule on this feature set, as discussed in Section 3.2.2.

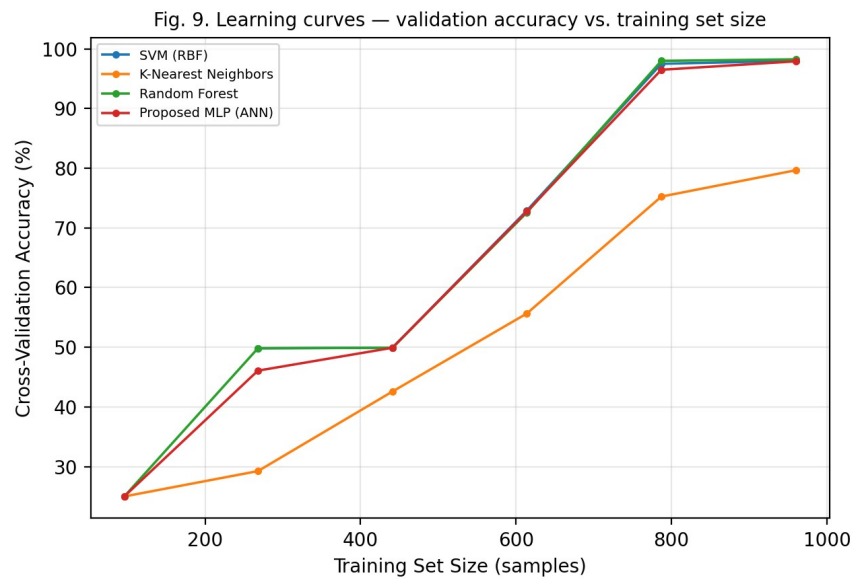


Fig. 9. Five-fold cross-validated learning curves showing validation accuracy as a function of training-set size for all four classifiers

5.10. Computational Cost Analysis

Table 9 and Figure 11 report training time and per-sample inference latency for each classifier, measured on a single CPU core as described in Section 4.9. The proposed ANN offers the lowest per-sample inference latency of the four classifiers (0.0013 ms/sample), owing to its fixed, small number of matrix-multiplication operations at inference time (484 trainable parameters in its tuned single-hidden-layer configuration), which is particularly relevant for high-sample-rate, real-time embedded deployment. SVM offers comparably low inference latency (0.0063 ms/sample) at very low training cost, though its inference cost scales with the number of support vectors retained from training, unlike the ANN's fixed parameter count. Random Forest, while statistically tied with the ANN on accuracy and the most robust to noise, incurs the highest training time and an inference latency roughly 53× that of the ANN, because each prediction requires traversing all 300 trees in the tuned ensemble. K-NN requires effectively no training cost, since it is an instance-based (lazy-learning) method, but incurs inference cost that grows with training-set size, since it must compute a distance to every stored training instance at prediction time.

Table 9. Computational cost comparison: training time, per-sample inference latency, and trainable parameter count.

Model	Training Time	Inference Time (per sample)	Trainable Parameters
SVM (RBF)	48.06 ms	0.0063 ms	n/a
K-Nearest Neighbors	1.52 ms	0.0152 ms	n/a
Random Forest	549.95 ms	0.0664 ms	n/a
Proposed MLP (ANN)	567.59 ms	0.0013 ms	484

Fig. 11. Computational cost comparison across models

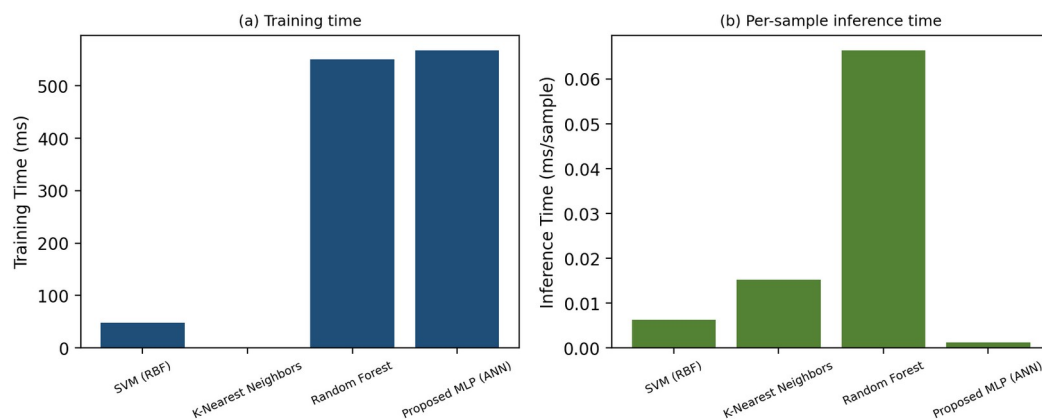


Fig. 11. Training time (a) and per-sample inference latency (b) comparison across the four evaluated classifiers

Taken together with the accuracy, statistical-significance, and noise-robustness findings above, this computational-cost analysis suggests that the proposed ANN offers the most favorable overall balance for real-time embedded deployment among the four classifiers evaluated: statistically indistinguishable accuracy from the best-performing alternatives, competitive noise robustness, and the lowest inference latency and most predictable (fixed) computational cost at prediction time, discussed further.

5.11. Comparison with Existing Literature

Table 8 situates the present study's results alongside a selection of the works reviewed in Section 2 that report quantitative accuracy figures. Direct numerical comparison across studies must be treated with caution, since sensing modality (current versus vibration versus operating parameters), fault taxonomy, dataset size, and evaluation protocol all differ; the comparison is included primarily to establish that the accuracy levels achieved in this study (upper-90s percent for SVM, Random Forest, and ANN) fall within the range reported by comparable feature-based classical-ML studies in the literature, rather than to claim outright superiority over any specific prior work.

Table 8. Qualitative comparison of the present study with selected works from the literature review (Section 2)

Study	Sensing Modality	Classifier(s)	Reported Accuracy
Kim et al. [2] (2023)	Vibration	SVM, MNN, CNN, GBM, XGBoost	High accuracy across models (per-model figures in source)
Pohakar et al. [5] (2025)	Operating parameters (V, I, speed)	RF, KNN, GBM, SVM, XGBoost+FIS	SVM mean accuracy > 90%
Abdulkareem et al. [6] (2025)	Voltage / current	ANN, Decision Tree, RF, k-NN	ANN approx. 91% after 40 epochs
Ullah et al. [7] (2025)	Vibration (MAFAULDA)	SVM, KNN, DNN+FFT	SVM 95.4%, KNN 92.8%, DNN+FFT 99.7%
Sobhi et al. [4] (2023)	Current / vibration (small motors)	Anomaly-detection ML pipeline	Effective for small-motor fleets (qualitative)
Proposed (this work)	Simulated stator current (MCSA-based)	SVM, k-NN, RF, ANN (proposed)	ANN and RF 98.33% (tuned, held-out test)

5.12. Discussion Summary

The results presented in this section support several consistent conclusions. First, a compact, ten-feature time/frequency representation is sufficient to separate the four modeled motor conditions with high accuracy (upper-90s percent) across three of the four evaluated classifier families, indicating that the discriminative information content of MCSA-derived statistical and spectral features is high relative to their computational cost, corroborating similar conclusions drawn from feature-based studies reviewed in Section 2. Second, classifier choice remains meaningful in one specific respect: k-NN significantly underperforms SVM, Random Forest, and the proposed ANN, which are themselves statistically indistinguishable from one another on accuracy, so that classifier selection among these three should be guided by secondary criteria this study finds noise robustness favors Random Forest, while inference latency and parameter-count predictability favor the proposed ANN. Third, the feature-ablation study identifies the sideband-energy ratio and harmonic ratio as carrying the large majority of the discriminative signal in this feature set, with the remaining eight features contributing comparatively modest, individually near-redundant information a finding with direct relevance to minimizing feature-computation cost in constrained deployments. Fourth, both noise robustness and data efficiency remain meaningfully imperfect even at the full dataset size and baseline noise level, indicating that further gains are plausible from larger, more diverse training data and from noise-robust feature or model design, both flagged as priorities for future work in Section 8.

5.13. Practical Recommendations

Based on the combined accuracy, statistical-significance, ablation, robustness, and computational-cost findings, three practical recommendations can be distilled for practitioners considering a similar MCSA-based, feature-driven fault-classification system. First, when measurement conditions are expected to be relatively clean and stable (for example, a well-instrumented laboratory or a low-noise industrial environment), the proposed ANN is an attractive default choice on the strength of its statistically top-tier accuracy combined with the lowest inference latency and smallest, most predictable memory footprint of the four classifiers evaluated. Second, when measurement noise is expected to be more variable or severe for instance, retrofit installations with lower-quality current transducers, or electrically noisy environments Random Forest's superior noise-robustness may outweigh its higher inference cost, particularly in applications where multiple motors are monitored from a shared, non-latency-critical edge gateway rather than an individually embedded microcontroller. Third, irrespective of which classifier is ultimately selected, the ablation results of Section 5.7 indicate that engineering effort is best concentrated on obtaining a clean, well-resolved frequency-domain estimate for the sideband-energy-ratio and harmonic-ratio features specifically, since these two features drive the large majority of achievable accuracy; conversely, k-NN is not recommended for this application under any of the tested conditions, given its significantly weaker and less noise-robust performance throughout Sections 5.3, 5.4, and 5.8.

6. Real-Time and Embedded Implementation Considerations

6.1. Embedded Deployment

The computational-cost results indicate that all four evaluated classifiers are, in principle, compatible with deployment on resource-constrained embedded hardware such as a digital signal processor (DSP), microcontroller, or field-programmable gate array (FPGA) integrated into a motor drive or a retrofit monitoring module, since even the most expensive model (Random Forest) requires well under one millisecond per prediction on a single, unaccelerated CPU core. This positions the proposed classifiers as a natural complement to the broader class of intelligent control architectures reviewed by Mishra, Mishra, and

Agarwal [27], who survey control paradigms ranging from classical PID and state-space methods to AI-driven approaches such as reinforcement learning and neuro-fuzzy control; in a fully integrated drive system, the fault-classification output of the present framework could plausibly feed such a higher-level intelligent controller to trigger derating, controlled shutdown, or maintenance scheduling in response to a detected fault. The dominant computational cost in a full embedded pipeline is more likely to be the fast Fourier transform (FFT) required for the frequency-domain features (sideband energy ratio, harmonic ratio, spectral centroid) than the classifier inference step itself; a 5,000-point FFT, computed once per one-second window, is comfortably within the throughput of low-cost embedded DSP cores and dedicated FFT hardware blocks available on many modern microcontrollers. The proposed ANN's fixed, small parameter count (484 parameters in its tuned configuration) is particularly favorable for deployment on memory-constrained devices, since its entire parameter set can be stored in a few kilobytes and inference reduces to two small matrix-vector multiplications and elementwise non-linearities, avoiding the branching, tree-traversal memory-access pattern required by Random Forest inference.

6.2. Latency and Throughput

For a condition-monitoring application evaluating one-second current windows, the effective diagnostic update rate is fundamentally limited by the window length itself (1 Hz in this study's configuration) rather than by classifier inference latency, which is several orders of magnitude faster. This suggests that shorter analysis windows could be explored to increase diagnostic update rate without classifier inference becoming a bottleneck, provided that feature quality (particularly frequency resolution for the sideband and harmonic features, which depends on window length) is not excessively degraded; this trade-off between window length, spectral resolution, and diagnostic update rate is a natural target for future work, noted in Section 8. For applications requiring simultaneous monitoring of many motors from a single edge gateway or industrial PC, the sub-millisecond per-sample inference times reported in Table 9 indicate that even the slowest evaluated classifier (Random Forest) could in principle support real-time monitoring of hundreds of motors from a single moderate-performance processing unit, well before inference latency becomes a practical constraint.

6.3. Integration with SCADA and Industrial IoT Infrastructure

Consistent with the cyber-physical-systems view of Industry 4.0-based manufacturing articulated by Lee, Bagheri, and Kao [23], in which physical assets are continuously mirrored by computational models and analytics to enable predictive rather than reactive operation, the proposed framework would sit downstream of existing current-sensing infrastructure already present for motor protection (e.g., overcurrent relays) or variable-frequency-drive control, with the feature-extraction and classification pipeline implemented either directly on drive firmware, on a dedicated edge-computing module, or on an Industrial Internet of Things (IIoT) gateway that aggregates data from multiple motors before forwarding diagnostic results to a supervisory control and data acquisition (SCADA) system or a plant-level condition-monitoring dashboard. This architectural pattern aligns closely with the layered CPS-IoT convergence framework reviewed by Mishra, Mishra, and Agarwal [26], who characterize the communication-protocol and real-time-data-processing layers linking edge-level sensing to plant- and enterprise-level analytics; the present framework's pipeline (Figure 12) maps naturally onto the sensing and edge-processing layers of such an architecture, with the classifier's predicted fault class and class probabilities constituting the payload passed upward to the network and application layers. Diagnostic outputs (predicted fault class and associated class probabilities from Section 5.6) map naturally onto standard SCADA alarm and trending mechanisms, and the class-probability outputs in particular allow configurable alarm thresholds per fault class, potentially prioritizing

early-warning sensitivity for costlier or more safety-critical fault types even at the cost of a higher false-alarm rate, consistent with the ROC/precision-recall threshold-tuning flexibility discussed in Section 5.6.

7. Limitations

The primary limitation of this study is its reliance on simulated rather than experimentally acquired current signatures. While the fault models were constructed from established MCSA fault-signature relationships (sideband frequencies, harmonic content, as summarized in Section 3.1) and deliberately randomized in severity, phase, and noise level to avoid a trivially separable synthetic classification problem, real motors exhibit additional sources of variability that were not modeled, including load transients and non-stationary operating conditions, supply-voltage imbalance and harmonic distortion from upstream power-quality issues, manufacturing tolerances and machine-to-machine variation, temperature-dependent parameter drift, and combined or simultaneous multi-fault conditions (addressed by only a small number of the reviewed studies, e.g. Pohakar et al. [5]). Validation on a public or laboratory-acquired current-signature dataset, ideally spanning multiple load conditions, multiple physical motors, and multiple fault severities, is a necessary next step before any deployment recommendation can be made with confidence.

A second limitation concerns the fault taxonomy: only three fault mechanisms plus the healthy condition were modeled, omitting air-gap eccentricity, broken end-ring segments, and external electrical faults such as single-phasing or voltage imbalance, all of which are documented in the reviewed literature (Section 2) as clinically relevant failure modes. A third limitation is that the noise-robustness and learning-curve experiments (Sections 5.8–5.9), while informative, each vary only one axis of difficulty (measurement noise or training-set size) independently; real-world performance depends on the joint interaction of data availability, noise, load variation, and fault severity, which was not exhaustively explored here. Finally, the ANN architecture evaluated is intentionally shallow; while this was found to be sufficient, and indeed optimal within the tested grid, for the ten-dimensional feature space used in this study, this finding should not be extrapolated to raw-waveform or deep-learning approaches (Section 2.4), which operate under a fundamentally different, higher-dimensional input representation and correspondingly different data requirements.

8. Conclusion and Future Work

This paper presented an extensive, reproducible machine learning framework for detecting and classifying bearing, rotor, and stator faults in three-phase induction motors from stator current signatures, and systematically compared four classifier families — SVM, k-NN, Random Forest, and a shallow ANN — under a common, hyperparameter-tuned, cross-validated evaluation protocol. The Random Forest and proposed ANN classifiers achieved the best held-out test accuracy (98.33% each), with SVM close behind (97.67%) and no statistically significant difference among these three (paired t-test, $p > 0.05$ in both comparisons); k-NN was significantly weaker (76.67%, $p < 0.001$ relative to the ANN). Beyond headline accuracy, a feature-ablation study identified the sideband-energy ratio and harmonic ratio as carrying the dominant share of diagnostic information in the ten-feature representation used, a noise-robustness sweep found Random Forest to degrade most gracefully under increasing measurement noise, learning curves indicated that all four classifiers remain data-limited even at the full dataset size, and a computational-cost analysis found the proposed ANN to offer the most favorable combination of accuracy, robustness, and inference latency for real-time embedded deployment among the classifiers evaluated.

Several directions for future work follow directly from the limitations identified in Section 7. First and most importantly, the framework should be validated on experimentally recorded current signatures, ideally spanning multiple physical motors, load conditions, and fault severities, to confirm that the accuracy, ablation, and robustness findings obtained here on simulated data transfer to real measurement conditions. Second, the fault taxonomy should be extended to include air-gap eccentricity, combined or simultaneous multi-fault conditions, and external electrical faults such as voltage imbalance, following the multi-fault methodology of studies such as Pohakar et al. [5]. Third, the reduced two-feature (sideband-energy ratio and harmonic-ratio) configuration suggested by the ablation study in Section 5.7 should be directly evaluated for its own accuracy and computational-cost trade-off, rather than only inferred from single-feature-removal results. Fourth, lightweight deep-learning architectures operating directly on raw or lightly processed current windows for example, one-dimensional convolutional networks, following the direction of Barrera-Llana et al. [13] and related work reviewed in Section 2.4 should be benchmarked against the feature-based classifiers evaluated here under the same tuned, cross-validated, robustness-aware protocol, to determine whether the additional representational flexibility of raw-waveform deep learning yields a meaningful accuracy or robustness advantage that justifies its higher data and computational requirements in this application. Finally, joint sensitivity experiments varying noise, load, and training-set size simultaneously, and a full hardware-in-the-loop evaluation on embedded target hardware, would further strengthen the practical deployment case for the framework proposed in this paper.

9. Appendix A: Nomenclature and Abbreviations

Table A1 summarizes the symbols and abbreviations used throughout this paper, for reference.

Symbol / Abbreviation	Definition
ANN	Artificial Neural Network
AP	Average Precision
AUC	Area Under the (ROC) Curve
BPFI / BPFO	Ball-Pass Frequency, Inner / Outer Race
BSF	Ball-Spin Frequency
CBM	Condition-Based Maintenance
CNN	Convolutional Neural Network
CV	Cross-Validation
DWT	Discrete Wavelet Transform
f_1	Fundamental (supply) frequency
FFT	Fast Fourier Transform
FTF	Fundamental Train (cage) Frequency

GBM	Gradient Boosting Machine
IIoT	Industrial Internet of Things
k-NN	k-Nearest Neighbors
MCSA	Motor Current Signature Analysis
MLP	Multilayer Perceptron
RBF	Radial Basis Function (SVM kernel)
RF	Random Forest
ROC	Receiver Operating Characteristic
s	Per-unit slip
SCADA	Supervisory Control and Data Acquisition
SVM	Support Vector Machine
THD	Total Harmonic Distortion
XGBoost	Extreme Gradient Boosting

10. Appendix B: Reproducibility Checklist

Following current best-practice recommendations for reporting machine-learning experiments in engineering research, Table B1 summarizes the reproducibility-relevant details of this study and the section in which each is documented, to facilitate independent replication or extension of the reported results.

Reproducibility Item	Status / Value
Random seed fixed	Yes (seed = 42), Section 4.9
Dataset generation procedure fully specified	Yes, Section 4.2
Feature definitions given as closed-form equations	Yes, Section 4.4
Train/test split ratio and stratification	75% / 25%, stratified, Section 4.7
Cross-validation protocol	5-fold stratified, Sections 4.7–4.8
Hyperparameter search space disclosed	Yes, Table 4 (Section 5.2)
Software environment specified	Python 3, NumPy, SciPy, scikit-learn, Matplotlib (Section 4.9)
Hardware environment specified	Single CPU core, no GPU acceleration (Section 4.9)

Statistical significance testing reported	Yes, paired t-test, Table 7 (Section 5.4)
Ablation study included	Yes, Table 5 (Section 5.7)
Robustness (noise) analysis included	Yes, Figure 10 (Section 5.8)
Data efficiency (learning curve) analysis included	Yes, Figure 9 (Section 5.9)
Computational cost (latency, parameters) reported	Yes, Table 9 (Section 5.10)
Real experimental validation	Not yet performed — identified as future work (Sections 7–8)

Author Contributions

Conceptualization, methodology, and simulation framework design: [Author 1]. Software implementation, feature engineering, and experiments (hyperparameter tuning, ablation, robustness, and computational-cost analysis): [Author 1]. Formal analysis and statistical testing: [Author 1]. Writing — original draft preparation: [Author 1]. Writing — review and editing: [Author 1, Author 2]. Supervision: [Author 2]. All authors have read and agreed to the submitted version of the manuscript.

Funding

This research received no external funding.

Data Availability Statement

The dataset-generation code, feature-extraction pipeline, and trained-model evaluation scripts used to produce the results reported in this paper are available from the corresponding author upon reasonable request. Because the primary dataset is synthetically generated according to the fully specified procedure of Section 4.2, it can be regenerated exactly given the stated random seed and parameterization, without requiring distribution of a static data file.

Conflicts of Interest

The authors declare no conflict of interest. The funders had no role in the design of the study; in the collection, analyses, or interpretation of data; in the writing of the manuscript; or in the decision to publish the results.

11. Acknowledgment

We are very much thankful to the authors of different publications as many new ideas are abstracted from them. Authors also express gratefulness to their colleagues and family members for their continuous help, inspirations, encouragement, and sacrifices without which this work could not be executed. Finally, the main target of this work will not be achieved unless it is used by research institutions, students, research scholars, and authors in their future works. The authors will remain ever grateful to Dr. Neelu Singh, Director, ICFRE Tropical Forest Research Institute, Jabalpur, Director, XLRI – Xavier School of Management, Jamshedpur & Principal Government Science College, Jabalpur who helped by giving constructive suggestions for this work.

The authors are also responsible for any possible errors and shortcomings, if any in the paper, despite the best attempt to make it immaculate.

References

- [1] P. Kumar and A. S. Hati, "Review on machine learning algorithm based fault detection in induction motors," *Archives of Computational Methods in Engineering*, vol. 28, no. 3, pp. 1929–1940, 2021, doi: 10.1007/s11831-020-09446-w.
- [2] M.-C. Kim, J.-H. Lee, D.-H. Wang, and I.-S. Lee, "Induction motor fault diagnosis using support vector machine, neural networks, and boosting methods," *Sensors*, vol. 23, no. 5, p. 2585, 2023, doi: 10.3390/s23052585.
- [3] M. Benninger, M. Liebschner, and C. Kreischer, "Fault detection of induction motors with combined modeling- and machine-learning-based framework," *Energies*, vol. 16, no. 8, p. 3429, 2023, doi: 10.3390/en16083429.
- [4] S. Sobhi, M. H. Reshadi, N. Zarft, A. Terheide, and S. Dick, "Condition monitoring and fault detection in small induction motors using machine learning algorithms," *Information*, vol. 14, no. 6, p. 329, 2023, doi: 10.3390/info14060329.
- [5] P. Pohakar, R. Gandhi, S. Hans, G. Sharma, and P. N. Bokoro, "Analysis of multiple faults in induction motor using machine learning techniques," *e-Prime — Advances in Electrical Engineering, Electronics and Energy*, vol. 12, art. 101007, 2025, doi: 10.1016/j.prime.2025.101007.
- [6] A. Abdulkareem, T. Anyim, O. Popoola, J. Abubakar, and A. Ayoade, "Prediction of induction motor faults using machine learning," *Heliyon*, vol. 11, art. e41493, 2025, doi: 10.1016/j.heliyon.2024.e41493.
- [7] I. Ullah, N. Khan, S. A. Memon, W.-G. Kim, J. Saleem, and S. Manzoor, "Vibration-based anomaly detection for induction motors using machine learning," *Sensors*, vol. 25, no. 3, p. 773, 2025, doi: 10.3390/s25030773.
- [8] J. L. Contreras-Hernandez, D. L. Almanza-Ojeda, S. Ledesma, A. Garcia-Perez, R. Castro-Sanchez, M. A. Gomez-Martinez, and M. A. Ibarra-Manzano, "Geometric analysis of signals for inference of multiple faults in induction motors," *Sensors*, vol. 22, no. 7, p. 2622, 2022, doi: 10.3390/s22072622.
- [9] M. E. H. Benbouzid, "A review of induction motors signature analysis as a medium for faults detection," *IEEE Transactions on Industrial Electronics*, vol. 47, no. 5, pp. 984–993, 2000.
- [10] O. E. Hassan, M. Amer, A. K. Abdelsalam, and B. W. Williams, "Induction motor broken rotor bar fault detection techniques based on fault signature analysis — a review," *IET Electric Power Applications*, vol. 12, no. 7, pp. 895–907, 2018.
- [11] A. Almounajjed, A. K. Sahoo, and M. K. Kumar, "Diagnosis of stator fault severity in induction motor based on discrete wavelet analysis," *Measurement*, vol. 182, art. 109780, 2021.
- [12] M. Zuhaib, F. A. Shaikh, W. Tanweer, A. M. Alnajim, S. Alyahya, S. Khan, M. Usman, M. Islam, and M. K. Hasan, "Faults feature extraction using discrete wavelet transform and artificial neural network for induction motor availability monitoring — Internet of Things enabled environment," *Energies*, vol. 15, no. 21, p. 7888, 2022, doi: 10.3390/en15217888.
- [13] K. Barrera-Llana, J. Burriel-Valencia, Á. Sapena-Bañó, and J. Martínez-Román, "A comparative analysis of deep learning convolutional neural network architectures for fault diagnosis of broken rotor bars in induction motors," *Sensors*, vol. 23, no. 19, p. 8196, 2023, doi: 10.3390/s23198196.
- [14] P. Gangsar and R. Tiwari, "Comparative investigation of vibration and current monitoring for prediction of mechanical and electrical faults in induction motor based on multiclass-support vector machine algorithms," *Mechanical Systems and Signal Processing*, vol. 94, pp. 464–481, 2017.
- [15] S. Misra, S. Kumar, S. Sayyad, A. Bongale, P. Jadhav, K. Kotecha, A. Abraham, and L. A. Gabralla, "Fault detection in induction motor using time domain and spectral imaging-based transfer learning approach on vibration data," *Sensors*, vol. 22, no. 21, p. 8210, 2022, doi: 10.3390/s22218210.
- [16] U. Ali, U. Ramzan, W. Ali, and K. A. Al-Jaafari, "An improved fault diagnosis strategy for induction motors using weighted probability ensemble deep learning," *IEEE Access*, 2025.
- [17] C. Cortes and V. Vapnik, "Support-vector networks," *Machine Learning*, vol. 20, no. 3, pp. 273–297, 1995.
- [18] T. Cover and P. Hart, "Nearest neighbor pattern classification," *IEEE Transactions on Information Theory*, vol. 13, no. 1, pp. 21–27, 1967.

- [19] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [20] D. E. Rumelhart, G. E. Hinton, and R. J. Williams, "Learning representations by back-propagating errors," *Nature*, vol. 323, pp. 533–536, 1986.
- [21] Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *Nature*, vol. 521, pp. 436–444, 2015.
- [22] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*, 2nd ed. New York, NY, USA: Springer, 2009.
- [23] J. Lee, B. Bagheri, and H.-A. Kao, "A cyber-physical systems architecture for Industry 4.0-based manufacturing systems," *Manufacturing Letters*, vol. 3, pp. 18–23, 2015.
- [24] K. Kudelina, T. Vaimann, B. Asad, A. Rassölkin, A. Kallaste, and G. Demidova, "Trends and challenges in intelligent condition monitoring of electrical machines using machine learning," *Applied Sciences*, vol. 11, no. 6, p. 2761, 2021, doi: 10.3390/app11062761.
- [25] K. Kudelina, B. Asad, T. Vaimann, A. Rassölkin, A. Kallaste, and A. Belahcen, "Methods of condition monitoring and fault detection for electrical machines," *Energies*, vol. 14, no. 22, p. 7459, 2021, doi: 10.3390/en14227459.
- [26] D. Mishra, R. K. Mishra, and R. Agarwal, "Cyber-Physical Systems and Internet of Things (IoT): Convergence, Architectures, and Engineering Applications," in *Advances in Engineering Science and Applications*, Jabalpur, India: Bhumi Publishing, 2025, pp. 17–43, ISBN: 978-93-48620-62-0.
- [27] R. K. Mishra, D. Mishra, and R. Agarwal, "Robotics and Intelligent Control Systems," in *Advances in Engineering Science: Theory, Technology, and Practice*, Nature Light Publications, 2025, doi: 10.5281/zenodo.16835664.
- [28] R. K. Mishra, D. Mishra, and R. Agarwal, *Big Data, Machine, and Deep Learning: Recent Progress, Key Applications, and Future Directions*, Munich, Germany: GRIN Publishing GmbH, 2025, ISBN (Book): 9783389122501, ISBN (eBook): 9783389122495.