

When Features Beat Noise: A Noise-Driven Hypothesis Testing Method for Feature Selection

Mousam Sinha¹, Tirtha Sarathi Ghosh², Ridam Pal³

¹Indian Statistical Institute (ISI), Kolkata, India

²Jadavpur University, Kolkata, India

³Indraprastha Institute of Information Technology (IIIT-Delhi), India

Abstract

Feature selection has remained a daunting challenge in machine learning and artificial intelligence, where increasingly complex, high-dimensional datasets demand principled strategies for isolating the most informative predictors. Despite widespread adoption, many established techniques such as Boruta, Recursive Feature Elimination and other wrapper based methods suffer from notable limitations; some incur substantial computational cost, while others offer no definite statistical driven stopping criteria or assesses the significance of their importance scores. A common heuristic approach introduces multiple random noise features and retains all predictors ranked above the strongest noise feature. Although intuitive, this strategy lacks theoretical justification and depends heavily on heuristics. This paper proposes a novel feature selection method that addresses these limitations. Our approach introduces multiple random noise features and evaluates each feature’s importance against the maximum importance value among these noise features incorporating a non-parametric bootstrap-based hypothesis testing framework to establish a solid theoretical foundation. We establish the conceptual soundness of our approach through statistical derivations that articulate the principles guiding the design of our algorithm. To evaluate its reliability, we generated simulated datasets under controlled statistical settings and benchmarked performance against Boruta and Knockoff-based methods, observing consistently stronger recovery of meaningful signal. As a demonstration of practical utility, we applied the technique across diverse real-world datasets, where it surpassed feature selection techniques including Boruta, RFE, and Extra Trees. Hence, the method emerges as a robust algorithm for principled feature selection, enabling the distillation of informative predictors that support reliable inference, enhanced predictive performance, and efficient computation.

Keywords: Feature selection; Multiple hypothesis testing; Bootstrap inference; Cross-validation; Boruta; Knockoff; Recursive feature elimination; Extra Trees, Hyperparameter tuning.

1 Introduction

With recent advancements in Machine learning, Deep learning, and Transformer architectures, various applications of high-dimensional data have grown rapidly, fundamentally transforming the landscape of classical predictive modeling [Borah et al. \(2024\)](#); [Rong et al. \(2019\)](#). Healthcare, bio-informatics, finance, marketing, and materials science now routinely generate structured datasets with thousands of features, capturing complex data patterns through diverse signals. Although rich with potential, the surge in dimensionality

also poses a critical challenge. It becomes essential to retain truly informative variables for downstream predictive tasks. High-dimensional data often contain redundant, irrelevant, or noisy variables, which can mask the underlying signal. This phenomenon makes it difficult for the machine learning models to capture the true data patterns leading to reduced predictive reliability. Isolating the relevant features would effectively enhance predictive performance, accelerate training, reduce overfitting, and improve the interpretability of complex models [Brahim Belhaouari et al. \(2023\)](#). Hence, feature selection is pivotal in developing robust and interpretable machine learning and deep learning models. While in deep learning, although representation learning inherently performs feature abstraction, implementing feature selection strategies can further lead to better model efficacy, faster convergence, reduced computational cost and enhanced model interpretability [Figueroa Barraza et al. \(2021\)](#).

Traditional approaches for feature selection are broadly categorized as filter [Sánchez-Marño et al. \(2007\)](#), wrapper [El Aboudi and Benhlina \(2016\)](#), or embedded methods [Wang et al. \(2015\)](#), each leveraging unique advantages [Jović et al. \(2015\)](#). Filter methods leverage statistical metrics to assess features independently, whereas wrapper methods evaluate subsets of features based on model performance. Embedded techniques incorporate feature selection during model training, often using regularization or tree-based strategies. Recent advancements such as SelectKBest, Recursive Feature Elimination (RFE) [Chen and Jeong \(2007\)](#), Lasso [Fonti and Belitser \(2017\)](#), and tree-based feature importance have shown higher efficacy in identifying relevant features. Among these, Boruta [Kursa et al. \(2010\)](#), a wrapper method built on random forests, stands out for its robustness in selecting all relevant features. The Extra-Trees classifier [Alfian et al. \(2022\)](#) has gained prominence for its ensemble-based importance estimation and resistance to overfitting.

Despite the proliferation of feature selection techniques in modern machine learning pipelines, several persistent limitations hinder their efficacy. Firstly, most commonly used algorithms such as Recursive Feature Elimination (RFE) lack an explicit statistical criterion to determine the optimal stopping point for feature retention. This absence of a quantitative threshold often leaves practitioners reliant on ad hoc or empirical decisions, which may compromise model performance. Secondly, when these selection procedures are coupled with computationally intensive models like tree-based models (e.g., Random Forest, XGBoost), the exhaustive need for hyperparameter tuning at each iteration makes the process expensive, particularly on large-scale or high-dimensional datasets. While methods like Boruta address the stopping rule by comparing features to their randomized “shadow” versions, this comes at the cost of doubling the feature space, exacerbating memory usage. Moreover, many existing approaches are fundamentally heuristic, empirical, and data-driven, lacking a principled mechanism to guarantee optimality in information retention or generalization. Perhaps, even after selecting a subset of features, current methods do not ensure these features do not contribute to model overfitting. These challenges highlight a pressing need for feature selection strategies that are both methodically robust and computationally efficient for distinguishing informative signals from noise while safeguarding against model overfitting.

2 Our Contribution

In this paper we tend to fill these gap by proposing a robust and interpretable feature selection technique that leverages model bootstrapping, noise calibration, and rigorous statistical testing to yield reliable predictive features with high generalization potential.

In Section 3, we introduce our proposed *Noise-Augmented Bootstrap Feature Selection (NABFS)* framework, which integrates controlled noise augmentation, resampling, and non-parametric statistical inference

to robustly identify informative features. The method begins by augmenting the real dataset with several artificially generated “noise” features drawn from a fixed random distribution, which act as a benchmark for detecting spurious importance. Using bootstrap resampling, the predictive model is repeatedly fitted to multiple resampled datasets, and feature importance are computed for both real and synthetic variables in each iteration. For every bootstrap replicate, the relative importance of each real feature is quantified against the maximum importance observed among the noise features. To assess whether a given feature consistently outperforms the noise benchmark, a one-sided Wilcoxon signed-rank test is performed across bootstrap iterations, yielding a p-value that captures the statistical evidence of informativeness. To control the overall false discovery rate, p-values are adjusted using the Holm–Bonferroni procedure, ensuring rigorous multiple testing correction. Finally, features with adjusted p-values less than the predefined significance threshold, α , are retained as significantly informative features. This approach not only enhances robustness against spurious correlations and overfitting but also provides a statistically grounded mechanism for distinguishing meaningful signals from noise within complex and high-dimensional data.

In Sections 5.1 and 5.2, we benchmark the proposed *Noise-Augmented Bootstrap Feature Selection (NABFS)* method against established approaches using both controlled simulation studies and real-world datasets. This dual evaluation framework enables a comprehensive assessment of the method’s effectiveness across diverse data settings, thereby demonstrating its robustness and generalizability.

In Section 6, we present a detailed discussion of our findings, highlighting the comparative performance of the proposed method against established feature selection techniques and outlining the future scope of this study. The results demonstrate that our approach consistently outperforms state-of-the-art methods across multiple evaluation settings, establishing a strong empirical benchmark. Comprehensive assessments across diverse machine learning models (logistic regression, random forest, gradient boosting, and extreme gradient boosting classifiers) reveal that models trained on the selected feature subsets achieve superior AUC scores compared to their full-feature counterparts, depicting the robustness of the proposed method in both linear and non-linear regimes. These findings reaffirm that even in an era dominated by complex architectures and data abundance, the principled selection of informative features remains indispensable for enhancing predictive accuracy, interpretability, and computational efficiency in artificial intelligent systems.

Hence, this paper provides a comprehensive theoretical foundation for the proposed method, supported by rigorous empirical validation. The effectiveness of the proposed algorithm is subsequently demonstrated through extensive evaluations on diverse datasets, where it is bench marked against state-of-the-art feature selection techniques to establish its robustness and practical utility.

3 Methods

3.1 Population Assumptions

We consider a population characterized by a response variable Y and p candidate features

$$f_1, f_2, \dots, f_p.$$

Each feature may or may not be statistically related to the response.

Our primary goal is to identify which features exhibit any form of statistical dependence with Y . For-

mally, for each feature f_j , we test the population-level hypothesis:

$$H_{0j} : Y \perp\!\!\!\perp f_j \quad \text{versus} \quad H_{1j} : Y \not\perp\!\!\!\perp f_j,$$

where $\perp\!\!\!\perp$ denotes statistical independence. Under the null, the feature f_j carries no information about Y .

To quantify this relationship, we define a *population-level importance measure*

$$\mathcal{I}(f_j, Y) = \text{a population measure of dependence between } f_j \text{ and } Y,$$

such as mutual information, correlation, regression coefficient magnitude, or feature importance metrics in tree based algorithms (e.g., XGB) like Gain, SHAP etc.

Equivalently, the hypotheses can be expressed as:

$$H_{0j} : \mathcal{I}(f_j, Y) = 0, \quad H_{1j} : \mathcal{I}(f_j, Y) > 0.$$

The measure $\mathcal{I}(f_j, Y)$ is assumed to be nonnegative, for instance when defined as an absolute correlation, mutual information, or coefficient magnitude. Hence, the alternative hypothesis corresponds to the presence of strictly positive population importance, indicating statistical dependence between f_j and Y .

3.2 Approximate solution

Because population importances $\mathcal{I}(f)$ are not directly observable, we construct an approximate solution based on resampling and empirical feature importance scores. The idea is to solve this hypothesis testing problem by repeatedly perturbing the data, estimating feature importances, and comparing candidate features against synthetic noise variables.

The procedure begins by augmenting the given data set with l artificial features drawn independently from a fixed distribution $\mathcal{D}_\varepsilon$. These noise variables play the role of a baseline: any real feature that cannot consistently outperform them should not be considered informative.

To approximate the distribution of importances, we employ the bootstrap. Specifically, we generate b bootstrap replicates of the augmented dataset, each of size n . For each replicate, we fit a predictive model, treated as a hyperparameter of the framework, and extract feature importance scores. Denote by $I_j^{(i)}(l)$ the importance of candidate feature f_j in the i -th bootstrap when l noise features are present, and by

$$M^{(i)}(l) = \max_{1 \leq k \leq l} I^{(i)}(\varepsilon_k; l)$$

the maximum importance among the l noise variables in the same replicate.

This construction yields a paired sequence of differences

$$D_i^{(j)}(l) = I_j^{(i)}(l) - M^{(i)}(l), \quad i = 1, \dots, b,$$

one per bootstrap replicate. For notational simplicity, we suppress the feature index j in $D_i^{(j)}(l)$ in what follows. We then test the one-sided alternative that the candidate feature tends to exceed the noise maximum

using the *Wilcoxon signed-rank (WSR)* statistic

$$T^+(l) = \sum_{i=1}^b R_i(l) \mathbf{1}\{D_i(l) > 0\},$$

where $R_i(l)$ is the rank of $|D_i(l)|$ among the nonzero $\{|D_1(l)|, \dots, |D_b(l)|\}$ (mid-ranks for ties; zero differences dropped). Large values of $T^+(l)$ indicate evidence in favor of the alternative H_{1j} . A one-sided p -value is obtained from the exact or large-sample WSR null (suppressed the feature index j).

Finally, because multiple features are tested simultaneously, p -values are adjusted using the Holm–Bonferroni procedure to control the family-wise error rate (FWER) at level α . Features that remain significant after adjustment are retained.

3.3 Algorithmic Summary

The algorithm below outlines the proposed Noise-Augmented Bootstrap Feature Selection (NABFS) method.

Require: Dataset (X, Y) with p candidate features;

l : number of synthetic noise features;

b : number of bootstrap replicates;

$\mathcal{M}$: choice of predictive model;

α : significance level (Type 1 error/ FWER control);

$\mathcal{D}_\varepsilon$: distribution of noise features (fixed).

Return : Selected feature set $\hat{\mathcal{S}} \subseteq \{f_1, \dots, f_p\}$.

1. Augment the dataset with l synthetic noise features sampled i.i.d. from $\mathcal{D}_\varepsilon$.
2. For each bootstrap replicate $i = 1, \dots, b$:
 - (a) Draw a bootstrap sample of size n from (X, Y) .
 - (b) Fit the chosen model $\mathcal{M}$ and compute feature importances.
 - (c) Record $I_j^{(i)}(l)$ for each candidate feature f_j and the maximum noise importance

$$M^{(i)}(l) = \max_{1 \leq k \leq l} I^{(i)}(\varepsilon_k; l).$$

- (d) Form the paired difference $D_i(l) = I_j^{(i)}(l) - M^{(i)}(l)$ for each f_j .
3. For each feature f_j , compute the one-sided Wilcoxon signed-rank statistic $T^+(l)$ from $\{D_i(l)\}_{i=1}^b$ and obtain a p -value.
4. Apply the Holm–Bonferroni procedure at level α to adjust p -values.
5. Return $\hat{\mathcal{S}} = \{f_j : \text{adjusted } p\text{-value} < \alpha\}$.

3.4 Practical Considerations

Several elements of the framework are treated as hyperparameters:

- **Number of noise features (l):** Larger l yields a stricter benchmark, reducing false positives but lowering power. Our theoretical analysis (Section 4) shows that both Type I error and power decrease monotonically with l under mild assumptions.
- **Number of bootstrap replicates (b):** Increasing b stabilizes the empirical distributions but raises

computational cost. In practice, b is chosen to balance stability with computational feasibility.

- **Choice of predictive model ($\mathcal{M}$):** The framework is model-agnostic: any learner that provides feature importance scores may be specified. Examples include tree-based models, linear models with coefficients as importance measures, or neural networks with attribution scores. In practice, one may select $\mathcal{M}$ from a candidate set by held-out predictive performance or stability of selections.
- **Significance level (α):** The target FWER level is a fixed input that determines the stringency of the Holm–Bonferroni correction. Smaller values of α reduce false discoveries but may also exclude weak signals.
- **Noise distribution ($\mathcal{D}_\varepsilon$):** Noise features are sampled independently from a fixed distribution. While Gaussian noise is common, other distributions (e.g., uniform) may be used. The choice of $\mathcal{D}_\varepsilon$ affects the scale of the benchmark but not the validity of the testing framework.

The procedure is approximate, relying on bootstrap distributions of empirical importances as proxies for unobservable population importances. Despite this approximation, as demonstrated in Section 5, the method provides a robust and practical tool for distinguishing informative features from noise. In applications, α is typically fixed at a conventional level (e.g., 0.05), and $\mathcal{D}_\varepsilon$ can be taken as standard Gaussian. The key tunable parameters are the number of noise features l and the underlying model $\mathcal{M}$; their choices may be guided by cross-validation, predictive performance on held-out data, and stability analyses of selected feature sets (see Section 5).

4 Theoretical Properties

The procedure of Section 2 provides an approximate solution to the population hypothesis testing problem of Section 2.1. In this section we establish its theoretical properties. We show that the number of synthetic noise features l governs a key trade-off between Type I error and statistical power. All proofs are deferred to the Appendix.

4.1 Type I Error and Power as Functions of l

We next study how the number of synthetic noise features, l , affects the statistical properties of the test. Recall that for each bootstrap replicate i , the paired difference is

$$D_i(l) = I_j^{(i)}(l) - M^{(i)}(l), \quad M^{(i)}(l) = \max_{1 \leq k \leq l} I^{(i)}(\varepsilon_k; l),$$

and the Wilcoxon signed-rank statistic is

$$T^+(l) = \sum_{i=1}^b R_i(l) \mathbf{1}\{D_i(l) > 0\},$$

where $R_i(l)$ is the rank of $|D_i(l)|$ among the nonzero absolute values.

Proposition 4.1 (Type I Error is Decreasing in l). *For any fixed feature under the null hypothesis, the Type I error probability decreases as l increases. If*

$$\alpha_j(l) = \mathbb{P}(T^+(l) \geq c_\alpha(b) \mid H_{0j} \text{ true}),$$

then

$$\alpha_j(l+1) \leq \alpha_j(l).$$

Intuition: Adding more noise features raises the replicate-wise maximum noise benchmark $M^{(i)}(l)$, which weakly lowers the paired differences $D_i(l)$; hence the WSR statistic $T^+(l)$ cannot increase, making false rejections less likely.

Proposition 4.2 (Power is Decreasing in l). *For any fixed feature under the alternative hypothesis, the power of the test decreases as l increases. If*

$$\beta_j(l) = \mathbb{P}(T^+(l) \geq c_\alpha(b) \mid H_{1j} \text{ true}),$$

then

$$\beta_j(l+1) \leq \beta_j(l).$$

Intuition. Increasing l makes the noise maximum stricter, so even truly informative features are less likely to consistently exceed the benchmark. Because $D_i(l)$ is non-increasing in l by Assumption (A1), the WSR statistic $T^+(l)$ is also pointwise non-increasing (see Appendix: WSR monotonicity lemmas), and thus the probability of crossing any fixed rejection threshold decreases.

4.2 Practical Implications

The two propositions show that the noise parameter l governs a fundamental trade-off: larger l reduces both Type I error and statistical power, producing a more conservative procedure; smaller l yields higher power but risks including spurious features. In practice, l should be treated as a hyperparameter and tuned to balance conservativeness with sensitivity to weak signals.

5 Data and experiments

5.1 Simulation Studies

We also synthesized data mimicing real use cases, conditioning boundary and edge cases. While simulating the data, we considered inclusion of factors such as sparsity, collinearity, and noise, which can allow for the systematic testing of the proposed approach across intricate settings. Hence, to systematically evaluate the performance of the proposed Noise-Augmented Bootstrap Feature Selection (NABFS) method, we conducted a comprehensive Monte Carlo simulation study against established feature selection techniques, which included Boruta [Kursa et al. \(2010\)](#) and Model-X Knockoffs [Candes et al. \(2018\)](#). The study was designed to assess the reliability and efficiency of each method under controlled yet diverse data-generating scenarios.

We simulated data from a logistic regression model with a compound symmetric correlation structure, a common framework for assessing feature selection performance in high-dimensional settings. Specifically, we designed a matrix $X \in \mathbb{R}^{n \times p}$, where each observation vector X_i was drawn as

$$X_i = \sqrt{1-\rho} Z_i + \sqrt{\rho} C,$$

with $Z_i \sim N(0, I_p)$ representing independent noise and $C \sim N(0, 1)$ denoting a shared latent factor introducing correlation among features. The true coefficient vector $\beta \in \mathbb{R}^p$ was defined such that the first k

features were true signals, each assigned a coefficient sampled uniformly from the interval $[-2, 2]$, while the remaining $p - k$ coefficients were set to zero. Binary outcomes were generated according to a logistic model,

$$\Pr(Y_i = 1 \mid X_i) = \frac{1}{1 + \exp(-X_i^\top \beta)}.$$

The experimental setup considered 50 candidate features ($p = 50$), of which 20 were true signals ($k = 20$), across sample sizes $n \in \{500, 1000, 3000\}$. The correlation parameter ρ varied across 12 levels $\{0, 0.01, 0.02, 0.05, 0.1, 0.2, 0.4, 0.5, 0.6, 0.8, 0.9, 1\}$, allowing systematic examination of the effect of feature correlation on selection stability. For NABFS, we varied the number of synthetic noise features ($l \in \{1, 2, 3, 4, 5, 6, 7\}$) to study its robustness across different augmentation intensities.

Each configuration was repeated 30 Monte Carlo replicates to ensure statistical reliability, with 3-fold cross-validation applied within each replicate for stability assessment.

For every combination of (n, ρ) , the data generation process, feature selection technique, and performance evaluation were repeated 30 times, and the averaged results across replicates were reported. The summary metrics with corresponding parameters $\{n, \rho, p, k\}$, which included average power, Type I error, and Jaccard index were compiled and have been reported in the result section.

This design enables a rigorous and reproducible comparison of NABFS with Boruta and Model-X Knockoffs across varying sample sizes, correlation structures, and noise conditions, providing a comprehensive understanding of the methods' stability, efficiency, and accuracy.

5.2 Real Data Application

We have tested our designed architecture for feature selection across diverse dataset to establish the generalizability of the technique. We catered datasets related to healthcare, genomics, proteomics, finance, and weather to effectively validate the robustness of our technique. In the previous section, we have demonstrated different scenario how we simulated the data based on various scenarios and tested the effectivity of the model.

Some of the real-world datasets related to healthcare were curated from the UCI Machine Learning Repository to understand the competence of feature selection technique. Clinical data describing the spectrum of complications following myocardial infarction, contributed by Golovenkin et al. [Golovenkin et al. \(2020\)](#), presented a rich domain-relevant attribute set suitable for rigorous medical data interrogation. The Parkinson's Disease Classification Dataset, made available by Sakar et al. [Sakar et al. \(2013\)](#), offers extensive biomedical voice measurements from individuals diagnosed with Parkinson's disease, enabling exploratory analysis of high-dimensional biological signals through advanced model selection methodologies. ShockModes dataset, introduced by Ridam Pal et al. [Pal et al. \(2022\)](#), is derived from the MIMIC-III intensive care unit database and comprises 17,294 ICU stays that were scored for the Shock Index (SI), a clinical metric used to prognosticate outcomes in critical care and emergency settings. The dataset identifies episodes of abnormal SI defined by sustained periods of elevated SI (> 0.7 for over 30 minutes) following a preceding window of normal SI, resulting in a curated cohort of 337 normal and 84 abnormal SI instances after exclusion criteria. The paper constituted of three cohorts, of which two cohorts, the multi-modal cohort and the uni-modal cohort (comprised of unstructured clinical notes embedding) were used for testing the effectivity of our method.

We even tested our dataset on molecular simulation data provided from the paper "Decrypting the mechanistic basis of CRISPR/Cas9 protein". This CRISPR/Cas9 study performed by Panda et al. [Panda and Ray](#)

(2024) used machine learning to classify metastable heteroduplex states in SpCas9 and FnCas9 complexes based on structural features derived from molecular dynamics simulations. The heteroduplex-state dataset contained 57,074 frames for SpCas9 and 52,080 for FnCas9, generated from long-timescale MD trajectories of both on-target and off-target systems. Each frame was labeled using a Markov State Model that identified discrete structural macrostates based on the dynamics of the RNA:DNA hybrid. For every frame, structural features were extracted, covering both RNA:DNA hybrid geometry (rotational and translational parameters for base pairs 17–20 relative to PAM, interfacial contacts, and interaction energy) and protein structural properties (backbone RMSD, radius of gyration, SASA, key angular descriptors, and inter-domain contacts). The dataset was organized to support classification using three input modalities: RNA:DNA hybrid-only, protein-only, and combined features, enabling systematic comparison of their contributions to distinguishing metastable states of the RNA-DNA heteroduplex.

For financial dataset, we catered data from Kaggle as an open resource. The Credit Card Fraud Detection dataset [Awoyemi et al. \(2017\)](#), curated by the Machine Learning Group of ULB and available via the Kaggle data platform, comprises anonymized transactions made by European cardholders over a two-day period in September 2013. The Loan Default Prediction Dataset [Robinson and Sindhwani \(2024\)](#), sourced from a Coursera challenge and made publicly available on Kaggle, includes detailed financial, demographic, and behavioral attributes describing borrowers’ profiles and loan histories.

The Student Performance Dataset, generated by Cortez [Silva \(2008\)](#), encompasses multidimensional academic, demographic, socio-economic, and institutional variables, providing a robust framework to benchmark feature selection approaches within educational research. The Airlines Customer Satisfaction dataset [Jethwa and Vora \(2024\)](#) comprises over 120,000 passenger records aimed at predicting overall satisfaction with airline services. It serves as a comprehensive benchmark for evaluating factors that influence passenger contentment and loyalty.

The collation of datasets from real-world setting both from open-source and private repository enables systematic and reproducible evaluation of our robust method (NABFS), offering insights into the robustness and adaptability of the proposed method under noise and artifact-prone scenarios. Table X summarizes key dataset characteristics which includes sample size, target variables, and core attributes while highlighting the heterogeneity and complexity of the data used in this study.

5.3 Metric Evaluation

The performance of the proposed method was systematically evaluated using a comprehensive set of metrics tailored for both simulated and real-world data. For simulation studies, model efficacy was assessed using *Power* [Cohen \(2013\)](#), *Type I Error* [Banerjee et al. \(2009\)](#), and the *Jaccard Index* [Real and Vargas \(1996\)](#), which collectively captures sensitivity, specificity, and selection stability. For real-world datasets, performance was further validated using the *F1-score* and *AUC-score*, reflecting the classification accuracy and discriminative capability of models trained on the selected feature subsets. A brief description of each metric is provided below.

- **Power:** Measures the proportion of true signal features correctly identified, indicating the model’s ability to detect relevant variables.
- **Type I error:** Represents the proportion of irrelevant (noise) features incorrectly selected, quantifying the model’s resistance to overfitting.

- **Jaccard index:** Evaluates the overlap between the selected and true feature sets, summarizing the overall selection accuracy and stability.
- **F1-score:** The harmonic mean of precision and recall, offering a balanced measure of classification performance, particularly in imbalanced datasets.
- **AUC-score:** The area under the ROC curve, providing a threshold-independent evaluation of the model’s discriminative strength.

Together, these metrics provide a holistic assessment of the proposed method, capturing its statistical reliability in controlled simulations and its predictive robustness on real-world datasets.

6 Results

Feature selection remains a daunting challenge in machine learning and deep learning, owing to the vast search space of potential predictors and the instability induced by correlated or noisy features. To address this, we developed the Noise-Augmented Bootstrap Feature Selection (NABFS), a principled approach designed to provide reliable feature selection in high-dimensional settings. The precedence of our methodology occurs in three stages. First, we examine limitations in existing probe-based feature selection method, particularly their susceptibility to correlation-induced instability. Next, we construct a refined theoretical formulation that establishes the statistical foundations of our approach. Finally, we evaluate the method across both simulated and real-world datasets to assess its empirical performance and generalizability. The assessment of NABFS was conducted in two broad segments, beginning with controlled simulation studies designed to characterize its behavior under specified conditions.

For assessing the performance of our proposed Noise-Augmented Bootstrap Feature Selection (NABFS) method, we designed an experimental setup of logistic regression framework with a compound symmetric correlation structure for data synthesis. To evaluate the performance of the proposed Noise-Augmented Bootstrap Feature Selection (NABFS) method, we developed a controlled simulation framework based on a logistic regression model with a compound-symmetric correlation structure for data synthesis. Two regimes were examined, one where the sample size exceeded the number of features ($n > p$) and another reflecting the high-dimensional setting ($p > n$). The simulations systematically varied the correlation parameter (ρ) to examine the influence of inter-feature dependence on selection performance. Each configuration was repeated across 30 Monte Carlo replicates to ensure statistical reliability, and results were averaged across runs. Comparative analyses were performed against two established baseline methods, Boruta and Model-X Knockoffs, across three sample sizes ($n = 300, 500$, and 800). The evaluation metrics included average Power, Type I Error, and Jaccard Index, as summarized in Figure 1.

Our experiments in the high-dimensional regime ($p > n$) were conducted with a fixed feature space of $p = 900$, comprising 700 true signal features and 200 synthetic noise features. Under this setting, for $n = 300$, NABFS demonstrates consistently higher average power across all correlation levels, particularly under low-to-moderate ρ values ($\rho < 0.5$). This indicates the method’s strong sensitivity in identifying true signal features even in limited data samples. Boruta shows gradual improvement with increasing correlation, while Knockoff remains largely conservative with negligible gain in power. The Type I error trends (Figure) reveal that NABFS maintains moderate false-positive rates, achieving a balanced trade-off between sensitivity and specificity. The Jaccard Index (Figure) further confirms NABFS’s robustness, maintaining



Figure 1: Performance of NABFS under the $p > n$ setting, evaluated using Power, Type I Error, and Jaccard Index across varying correlation levels.

superior overlap between the selected and true feature sets compared to both baselines. As the sample size increases to $n = 500$, the advantage of NABFS becomes more pronounced. The average power curve stabilizes at higher levels, reflecting improved statistical efficiency with larger data. Both Boruta and Knockoff show marginal improvement; however, Boruta continues to exhibit inflated Type I error, while Knockoff remains conservative and fails to recover true features effectively. NABFS sustains its equilibrium between power and false discovery, validating its scalability across varying data regimes. The Jaccard Index improves across all methods, but NABFS continues to achieve the highest stability and selection accuracy. For $n = 800$, NABFS exhibits near-saturation in power across all correlation levels, indicating that the majority of true informative features are consistently identified. Type I error remains well-controlled, with minimal increase despite higher data dimensionality. In contrast, Boruta’s variability widens under strong correlations ($\rho > 0.6$), while Knockoff continues to underperform in terms of power. The Jaccard Index trends reaffirm NABFS’s robustness and generalizability, maintaining the highest alignment with the ground-truth feature set across all conditions. Overall, across all sample sizes and correlation levels, NABFS achieves a robust balance between sensitivity and specificity, consistently outperforming existing feature selection methods. Its stable performance with increasing sample size emphasizes upon scalability and robustness, validating NABFS as a reliable feature selection technique for high-dimensional data with correlated and noisy features.

To evaluate NABFS in the $n > p$ setting, we considered a simulation design with $p = 50$ candidate features, of which $k = 20$ were true signals, and sample sizes varied across $n \in \{500, 1000, 3000\}$. This configuration enabled a controlled examination of how increasing data availability influences the performance and stability of feature selection technique. As shown in Figure 2, NABFS exhibits a consistent increase in average power with growing sample size, indicating its ability to accurately recover true signal features with data abundance. At $n = 500$, NABFS achieves moderate power (≈ 0.8 at $\rho = 0$) with controlled Type I error (< 0.01), maintaining balanced sensitivity and specificity. Although mild degradation

in power is observed at higher correlation levels ($\rho > 0.6$), the method remains stable, as reflected by the gradual decline in the Jaccard Index, demonstrating robustness against multicollinearity. When the sample size increases to $n = 1000$, NABFS's performance markedly improves, with average power exceeding 0.85 at low correlation and maintaining higher stability across ρ values. Type I error remains well-regulated for weak to moderate correlations ($\rho < 0.3$), suggestive of less noise retention even as feature dependencies increases. The Jaccard Index trends mirror the power trajectory, confirming that the selected feature sets increasingly align with the true informative subset as sample size grows. At the largest sample size, $n = 3000$, NABFS achieves near-saturated power across low-to-moderate correlations ($\rho < 0.5$), approaching perfect distillation of true signal features (Power ≈ 0.9 and Jaccard ≈ 0.85 at $\rho = 0$). Even under stronger correlation regimes, NABFS maintains competitive performance with modest Type I error inflation, highlighting its resilience in scenarios when correlated features are more. These results collectively affirm the robustness of NABFS in scenarios where data abundance enhances feature retention capacity. Overall, across all configurations, NABFS demonstrates a desirable trade-off between sensitivity and specificity, with consistent improvement in selection stability as data samples (n) increases. The convergence of Power and Jaccard Index values in the large-sample regime substantiates the theoretical underpinnings of NABFS, validating the efficacy of our technique.

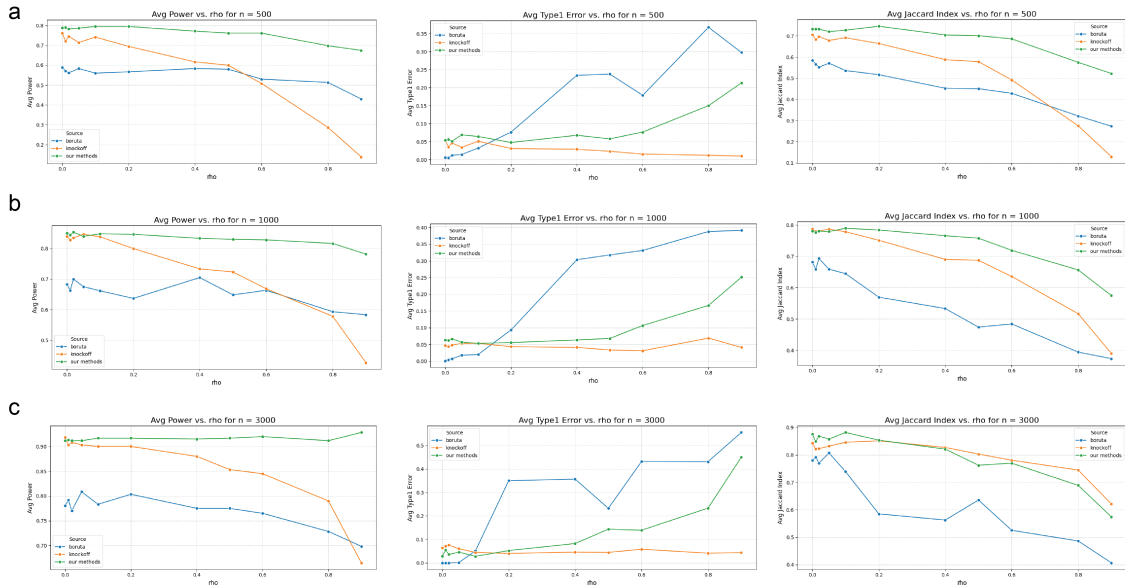


Figure 2: Performance of NABFS under the $n > p$ setting, evaluated using Power, Type I Error, and Jaccard Index across varying correlation levels.

Building on the extensive simulation studies that demonstrated the effectiveness of the proposed Noise-Augmented Bootstrap Feature Selection (NABFS) method, we further evaluated its practical utility across a diverse collection of real-world datasets (healthcare, finance and biological sequences). This empirical assessment was designed to test whether the advantages observed in controlled synthetic settings are also translated to heterogeneous, noisy and domain-specific datasets. NABFS was systematically benchmarked against established feature selection techniques, including Boruta, Recursive Feature Elimination (RFE) and Extra Trees–based importance scoring, using F1-score and Area Under the ROC Curve (AUC) as evaluation metrics. Hyper-parameter selection for our feature selection technique was based upon AUC-based perfor-

mance trends obtained across multiple noise intensities for both the model base learners (Logistic regression and XGBoost). For each base learner (Logistic and Xgboost), AUC values were evaluated across multiple noise intensities, and the final combination model–noise configuration was chosen based on the highest observed AUC, enabling principled, data-driven selection rather than heuristic tuning. The logistic regression–based model used a significance threshold of $\alpha = 0.01$, whereas the XGBoost-based model adopted $\alpha = 0.05$, reflecting differing sensitivity to linear and non-linear signal patterns.

Dataset	Our Technique (NABFS)			Boruta			RFE			Extra Tree Classifier		
	% Retained	F1	AUC	% Retained	F1	AUC	% Retained	F1	AUC	% Retained	F1	AUC
ShockModes (AdaBoost + W2V + Doc2Vec)	15	0.744	0.596	5	0.733	0.622	5	0.723	0.636	16	0.740	0.650
ShockModes (GBoost + W2V + Doc2Vec)	9	0.756	0.656	5	0.772	0.654	5	0.745	0.673	16	0.75	0.63
Parkinson’s Disease	12	0.711	0.827	16	0.707	0.824	16	0.695	0.823	–	–	–
Credit Card	59	0.999	0.968	100	0.999	0.967	100	0.999	0.967	–	–	–
Loan Default	75	0.838	0.748	46	0.839	0.744	46	0.839	0.745	–	–	–
Airline	100	0.814	0.905	35	0.800	0.873	35	0.800	0.873	–	–	–
Myocardial	19	0.863	0.737	2	0.857	0.685	2	0.855	0.617	–	–	–
Fn-Cas9	83	0.868	–	100	0.868	–	100	0.868	–	–	–	–
Sp-Cas9	100	0.947	–	100	0.947	–	100	0.947	–	–	–	–

Table 1: Performance comparison of NABFS against baseline feature selection methods across datasets for classification task. NABFS consistently achieves competitive predictive performance while substantially reducing model complexity across heterogeneous datasets.

Across healthcare and clinical datasets, NABFS consistently demonstrated strong discriminative performance while retaining compact and interpretable feature sets. In the multimodal ShockModes dataset, which integrates structured clinical variables with unstructured textual representations, NABFS exhibited robust performance across both model pipelines. For the AdaBoost + Word2Vec + Doc2Vec configuration, NABFS retained 15% of the original features while achieving an F1 score of 0.744 and an AUC of 0.596, closely approximating the performance of the Extra Trees baseline despite operating over a substantially reduced feature space (Table 1). These results highlight NABFS’s capacity to extract clinically relevant signal from high-dimensional, heterogeneous healthcare data.

Similarly, in the Parkinson’s disease classification task, NABFS achieved an AUC of 0.827 while retaining only 12% of the available features, outperforming both Boruta and RFE, which required larger subsets to achieve comparable performance. In the Myocardial dataset, NABFS attained an F1 score of 0.863 with an AUC of 0.737, exceeding comparator methods while maintaining substantial feature compression (Table 1).

In financial risk modelling tasks, NABFS preserved near-optimal predictive performance while improving model compactness. For the Credit Card fraud detection dataset, NABFS achieved an F1 score of 0.999 and an AUC of 0.968 while retaining 59% of the features, matching the near-ceiling performance of Boruta and RFE, which utilised the full feature set. In the Loan Default dataset, NABFS reached an F1 score of 0.838 and an AUC of 0.748 while retaining 75% of the features, delivering performance comparable to or exceeding baseline methods operating with larger or equivalent feature subsets (Table 1). These findings demonstrate that NABFS sustains strong discriminative capacity in high-stakes financial domains while reducing feature redundancy.

For biological sequence prediction tasks, including Fn-Cas9 and Sp-Cas9 activity modelling, NABFS achieved the strongest predictive performance among all evaluated methods. On the Fn-Cas9 dataset,

Dataset	Our Technique (NABFS)			Boruta			RFE		
	Retained Features (%)	RMSE	R^2	Retained Features (%)	RMSE	R^2	Retained Features (%)	RMSE	R^2
Student Performance	74%	3.656	0.099	13%	3.643	0.106	13%	3.691	0.082

Table 2: Performance comparison of NABFS against baseline feature selection methods across datasets for regression tasks. NABFS maintains competitive predictive accuracy with substantially different feature retention behaviour relative to Boruta and RFE.

NABFS obtained an F1 score of 0.868 while retaining 83% of features, and on the Sp-Cas9 dataset, it achieved an F1 score of 0.947 with full feature retention, matching or surpassing filter-based and wrapper-based comparators (Table 1). These results indicate that NABFS remains robust in domains characterised by highly structured and correlated biological signals, where preservation of subtle sequence-derived patterns is critical.

Finally, across other structured tabular datasets, NABFS maintained a favourable balance between predictive performance and feature efficiency. In the Airline delay dataset, NABFS achieved an F1 score of 0.814 and an AUC of 0.905, outperforming Boruta and RFE in terms of discrimination while retaining the full feature set (Table 1). For regression-based evaluation, we assessed NABFS on the Student Performance dataset, where the objective was to predict continuous academic outcomes. NABFS retained 74% of the available features and achieved an RMSE of 3.656 with an R^2 of 0.099, demonstrating performance comparable to Boruta and RFE despite markedly different feature selection behaviour (Table 2). Notably, while Boruta and RFE operated on substantially narrower feature subsets, their predictive gains remained marginal, indicating that NABFS preserves a broader set of informative attributes without introducing noise or instability. These findings suggest that NABFS extends effectively beyond classification paradigms, offering consistent performance in continuous outcome modelling tasks where nuanced feature interactions are critical. Collectively, these results demonstrate that NABFS effectively distils informative representations across heterogeneous data modalities, particularly in high-dimensional and clinically grounded settings, while preserving or improving predictive accuracy.

Together, these results demonstrate that the advantages of NABFS are not confined to controlled simulation regimes. Rather, the method generalizes effectively to real-world, multimodal and high-dimensional datasets maintaining competitive or superior predictive performance while offering more parsimonious, stable and interpretable feature sets. This combination of statistical rigor and practical efficiency positions NABFS as a robust and domain-agnostic framework for modern feature selection challenges.

7 Discussion & Conclusion

This study introduces Noise-Augmented Bootstrap Feature Selection (NABFS), a statistically principled approach for feature selection by incorporating formal hypothesis testing into the selection process. In contrast to heuristic probe-based methods that rely on a single noise comparison or empirically tuned thresholds, our method formulates feature relevance as a formal hypothesis-testing problem, leveraging bootstrap-stabilized importance estimates and a non-parametric Wilcoxon signed-rank test with Holm–Bonferroni correction. Such a foundation distinguishes NABFS from heuristic feature selection methods, providing greater confidence that the selected variables are not false discoveries but genuine predictors supported by evidence.

In both controlled simulations and diverse real-world datasets, NABFS demonstrated strong performance relative to established feature selection techniques. In simulation studies, NABFS consistently

achieved higher statistical power in detecting true relevant features compared to alternatives like Boruta and Model-X knockoffs, while maintaining a controlled false positive rate. Notably, even as inter-feature correlation increased, NABFS balanced sensitivity and specificity better than the alternatives, which tended to either suffer inflated false discoveries (in the case of Boruta) or remain overly conservative (in the case of knockoffs) with negligible gains in power. The selected feature sets from NABFS also showed greater stability and overlap with the true feature set than those of baseline methods, depicting its reliability in identifying meaningful signals without being misled by spurious correlations. This robust performance in simulations established a strong empirical benchmark for our technique. The efficacy of our method was translated into real-world applications, where NABFS delivered competitive or superior results using substantially reduced and more informative feature subsets. Our method consistently outperformed or achieved comparable predictive performance in comparison to established feature selection techniques (including Boruta, recursive feature elimination and Model-X knockoffs). Together, these results asserted that the benefits of NABFS are not just confined to simulations, rather the method generalizes effectively to real-world, multimodal and high-dimensional datasets, achieving competitive predictive performance while yielding more compact and interpretable feature sets. In fact, models trained on the NABFS-selected features often achieved higher accuracy (AUC) than models using the full original feature set.

A key strength of NABFS lies in its ability to explicitly manage the trade-off between statistical power and Type-I error. By varying the number of introduced noise features (m), practitioners can systematically control the statistical stringency of the selection criteria. Increasing the size of the noise feature pool imposes a more stringent empirical null threshold that true signal features must surpass, thereby reducing susceptibility to false discoveries at the expense of statistical power. Conversely, using fewer noise probes makes the test more permissive, increasing sensitivity to weak signals but potentially allowing more false positives. In essence, the method provides a tunable parameter m to attain balance between Power and Type I errors. Empirically, we observed that power is a decreasing function of m , meaning one can calibrate m to determine an optimal trade-off that suits the problem’s tolerance for false discoveries. This capability to regulate the error trade-off is a distinctive advantage as it offers a principled means to maintain high discovery power while keeping false discoveries in check. NABFS is also notable for its computational efficiency and model-agnostic flexibility. Despite the added overhead of bootstrap resampling and noise augmentation, the method remains tractable even in high-dimensional settings. Our observation states that NABFS often required fewer model training iterations than exhaustive approaches like RFE (which retrains models for many subset sizes), and also it does not dramatically inflate memory usage like methods that double the feature space (e.g. naive implementations of Boruta, Knocoff). Furthermore, unlike certain feature selection techniques tied to specific model families (for instance, Boruta’s dependence on random forests), NABFS can be paired with any predictive modeling algorithm without modification. The model-agnostic nature of NABFS was demonstrated by its seamless integration with both linear base learners (logistic regression) and non-linear ensemble base learners (including XGBoost and Random Forest), relying solely on each learner’s feature importance mechanism without requiring algorithm-specific modifications. In summary, the proposed method emerges as a robust yet efficient method for principled feature selection, enabling the distillation of informative predictors that support reliable inference, improved model performance, and reduced computational cost.

Despite its empirical strengths, NABFS has some limitations that warrant careful consideration. First, the method relies on a base learner to estimate feature importance and is therefore sensitive to the stability

and inductive biases of the underlying model. Second, the proposed procedure constitutes an approximate solution to the formal problem outlined in Section 3. While we attempted to control the family-wise error rate (FWER) through multiple hypothesis testing (via Holm–Bonferroni-type corrections), deriving strict theoretical guarantees of FWER control under our construction remains a future scope. The method further operates under a set of modelling assumptions described in Appendix Section A. If these assumptions are violated, there is no formal guarantee that statistical power or Type I error will behave as a monotonic function of the number of injected noise features (m). However, we provided intuition and empirical evidence suggesting that these assumptions are not overly restrictive in practical settings. Moreover, the current implementation injects synthetic noise drawn from a fixed Gaussian distribution, $\mathcal{N}(0, 0.01)$, which constitutes an additional tunable hyperparameter. Although fixed in this study for reproducibility, practitioners may select alternative noise distributions and variances more appropriate to their domain and requirement of the data. Like other importance-based selection frameworks, NABFS can be challenged by highly collinear or redundant predictors. In such settings, feature importance can be diluted across correlated variables, leading the method to preferentially select dominant representatives while discarding redundant yet meaningful features. Similarly, excessively stringent selection thresholds arising either from overly conservative significance levels or large values of m can suppress weak but genuinely informative signals by elevating the synthetic-noise baseline. While we provide principled default settings, all hyperparameters (bootstrap iterations, noise magnitude, m , and significance thresholds) remain user-tunable and should be optimised for domain-specific deployments. Finally, the current method is designed primarily to control the family-wise error rate. Extending the method to support explicit false discovery rate (FDR) control represents an important direction for future work. Collectively, these limitations highlight that while NABFS offers a practical robust approach, its theoretical guarantees remain approximate, which in turn requires careful tuning and domain expertise for optimal deployment.

As future scope, there are several promising directions for extending our research. First, while we use Holm–Bonferroni in this paper, it would be valuable to develop stronger theoretical guarantees for FWER control under the effect of bootstrapping. This requires new theory, since bootstrap samples are not fully independent, and exploring this formally would make the method more rigorous. Another direction can be to improve the efficacy of NABFS when features are highly correlated. Possible extensions include grouping correlated variables, using conditional importance scores, or adding a special step to avoid losing meaningful but redundant features. Future versions of NABFS could also use adaptive noise, where the scale or distribution of the injected noise automatically adjusts to the dataset instead of being fixed. This would make the method easier to apply across diverse domains of dataset with heterogeneous statistical characteristics.

In summary, NABFS offers a statistically grounded alternative to traditional feature-selection methods by combining noise augmentation, bootstrap resampling, and non-parametric testing. This design allows the method to reliably distinguish meaningful features from noise while keeping the selection process transparent and hyper-tunable. Across simulated and real-world datasets, NABFS selected compact feature sets without effecting the predictive performance, demonstrating its feasibility in settings where interpretability and stability are important. Future work may explore stronger theoretical guarantees, improved handling of correlated features, adaptive noise strategies, and extensions to settings such as deep learning or unsupervised learning. Together, these directions will make NABFS even more flexible and scalable. Overall, our method delivers a robust, statistically principled and computationally efficient approach to feature selection,

well suited for modern machine learning workflows.

References

- Ganjar Alfian, Muhammad Syafrudin, Imam Fahrurrozi, Norma Latif Fitriyani, Fransiskus Tatas Dwi Atmaji, Tri Widodo, Nurul Bahiyah, Filip Benes, and Jongtae Rhee. Predicting breast cancer from risk factors using svm and extra-trees-based feature selection method. *Computers*, 11(9):136, 2022.
- John O Awoyemi, Adebayo O Adetunmbi, and Samuel A Oluwadare. Credit card fraud detection using machine learning techniques: A comparative analysis. In *2017 international conference on computing networking and informatics (ICCNi)*, pages 1–9. IEEE, 2017.
- Amitav Banerjee, Umesh B Chitnis, Sudhir L Jadhav, Jitendra Shyamsundar Bhawalkar, and Suprakash Chaudhury. Hypothesis testing, type i and type ii errors. *Industrial psychiatry journal*, 18(2):127–131, 2009.
- Kasmika Borah, Himanish Shekhar Das, Soumita Seth, Koushik Mallick, Zubair Rahaman, and Saurav Mallik. A review on advancements in feature selection and feature extraction for high-dimensional ngs data analysis. *Functional & Integrative Genomics*, 24(5):139, 2024.
- Samir Brahim Belhaouari, Mohammed Bilal Shakeel, Aiman Erbad, Zarina Oflaz, and Khelil Kassoul. Bird’s eye view feature selection for high-dimensional data. *Scientific Reports*, 13(1):13303, 2023.
- Emmanuel Candes, Yingying Fan, Lucas Janson, and Jinchi Lv. Panning for gold: ‘model-x’ knockoffs for high dimensional controlled variable selection. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 80(3):551–577, 2018.
- Xue-wen Chen and Jong Cheol Jeong. Enhanced recursive feature elimination. In *Sixth international conference on machine learning and applications (ICMLA 2007)*, pages 429–435. IEEE, 2007.
- Jacob Cohen. *Statistical power analysis for the behavioral sciences*. routledge, 2013.
- Naoual El Aboudi and Laila Benhlila. Review on wrapper feature selection approaches. In *2016 international conference on engineering & MIS (ICEMIS)*, pages 1–5. IEEE, 2016.
- Joaquín Figueroa Barraza, Enrique López Droguett, and Marcelo Ramos Martins. Towards interpretable deep learning: a feature selection framework for prognostics and health management using deep neural networks. *Sensors*, 21(17):5888, 2021.
- Valeria Fonti and Eduard Belitser. Feature selection using lasso. *VU Amsterdam research paper in business analytics*, 30:1–25, 2017.
- Sergey E Golovenkin, Jonathan Bac, Alexander Chervov, Evgeny M Mirkes, Yuliya V Orlova, Emmanuel Barillot, Alexander N Gorban, and Andrei Zinovyev. Trajectories, bifurcations, and pseudo-time in large clinical datasets: applications to myocardial infarction and diabetes data. *GigaScience*, 9(11):giaa128, 2020.
- Sejal Jethwa and Neha Vora. A comprehensive analysis of airline passenger’s satisfaction through classification model. *Journal of Trends in Computer Science and Smart Technology*, 6(4):324–337, 2024.

- Alan Jović, Karla Brkić, and Nikola Bogunović. A review of feature selection methods with applications. In *2015 38th international convention on information and communication technology, electronics and microelectronics (MIPRO)*, pages 1200–1205. Ieee, 2015.
- Miron B Kursa, Aleksander Jankowski, and Witold R Rudnicki. Boruta—a system for feature selection. *Fundamenta informaticae*, 101(4):271–285, 2010.
- Ridam Pal, Shaswat Patel, Akshala Bhatnagar, Hardik Garg, Pradeep Singh, Ritesh Singh Soun, Aditya Agarwal, Aditya Nagori, Ashish Khanna, Rakesh Lodha, et al. Shockmodes: A multimodal model for prognosticating intensive care outcomes from physician notes and vitals. *medRxiv*, pages 2022–12, 2022.
- Gayatri Panda and Arjun Ray. Deciphering cas9 specificity: Role of domain dynamics and rna: Dna hybrid interactions revealed through machine learning and accelerated molecular simulations. *International Journal of Biological Macromolecules*, 283:137835, 2024.
- Raimundo Real and Juan M Vargas. The probabilistic basis of jaccard’s index of similarity. *Systematic biology*, 45(3):380–385, 1996.
- Natasha Robinson and Nidhi Sindhwani. Loan default prediction using machine learning. In *2024 11th International Conference on Reliability, Infocom Technologies and Optimization (Trends and Future Directions)(ICRITO)*, pages 1–5. IEEE, 2024.
- Miao Rong, Dunwei Gong, and Xiaozhi Gao. Feature selection and its use in big data: challenges, methods, and trends. *Ieee Access*, 7:19709–19725, 2019.
- Betul Erdogan Sakar, M Erdem Isenkul, C Okan Sakar, Ahmet Sertbas, Fikret Gorgen, Sakir Delil, Hulya Apaydin, and Olcay Kursun. Collection and analysis of a parkinson speech dataset with multiple types of sound recordings. *IEEE journal of biomedical and health informatics*, 17(4):828–834, 2013.
- Noelia Sánchez-Marroño, Amparo Alonso-Betanzos, and María Tombilla-Sanromán. Filter methods for feature selection—a comparative study. In *International conference on intelligent data engineering and automated learning*, pages 178–187. Springer, 2007.
- Alice Silva. Using data mining to predict secondary school student performance. *Journal Name Placeholder*, 2008.
- Suhang Wang, Jiliang Tang, and Huan Liu. Embedded unsupervised feature selection. In *Proceedings of the AAAI conference on artificial intelligence*, volume 29, 2015.

Appendix

A Proofs of Theoretical Results

Let f_j be a candidate feature. For bootstrap replicate $i = 1, \dots, b$, define:

- $I_j^{(i)}(l)$: the importance score of feature f_j when evaluated with l synthetic noise features included;

- $I^{(i)}(\varepsilon_k; l)$: the importance score of the k -th synthetic noise feature in replicate i when l noise features are included;
- $M^{(i)}(l) = \max_{1 \leq k \leq l} I^{(i)}(\varepsilon_k; l)$: the maximum noise importance at level l in replicate i ;
- $D_i(l) = I_j^{(i)}(l) - M^{(i)}(l)$: the paired difference between the signal feature and the maximum noise feature.

Let $T^+(l)$ be the one-sided Wilcoxon signed-rank statistic computed from $\{D_i(l)\}_{i=1}^b$, and let $c_\alpha(b)$ be the Wilcoxon critical value at level α , which depends only on b . Assume:

- (A1) (*Monotone differences*) $D_i(l+1) \leq D_i(l)$ for all i (e.g., it suffices that $I_j^{(i)}(l+1) \leq I_j^{(i)}(l)$ and $M^{(i)}(l+1) \geq M^{(i)}(l)$ pointwise).
- (A2) (*Continuity*) Feature importances are continuous (ties occur with probability zero; in practice, ties use mid-ranks and zeros are dropped, with b interpreted as the effective number of nonzero pairs).
- (A3) The paired differences $\{D_i(l)\}_{i=1}^b$ obtained from bootstrap replicates are treated as approximately independent draws from the sampling distribution of the importance statistic.

Assumption (A1) is motivated by the sequential nature of the construction. At step $l+1$, the data used at level l are held fixed while an additional independent noise feature is appended. Consequently, the maximum noise importance $M^{(i)}(l)$ tends to increase (or remain stable) with l , as it is computed over an expanding set of noise variables. If total importance is bounded, without loss of generality, normalized such that $\sum_j I_j^{(i)}(l) = 1$, then an increase in the importance of the newly added noise feature intuitively implies that the relative importance of some existing features is likely to decrease or remain unchanged. Hence, in expectation, the sequence of differences $D_i(l) = I_j^{(i)}(l) - M^{(i)}(l)$ can be viewed as nonincreasing in l , providing an intuitive justification for (A1).

Lemma A.1 (Wilcoxon Signed Rank (WSR) decomposition). *Let $\{D_i(l)\}_{i=1}^b$ denote the signed paired differences at level l , and let $T^+(l)$ denote the Wilcoxon positive rank sum statistic computed from the absolute values $\{|D_1(l)|, \dots, |D_b(l)|\}$, after discarding zeros and assuming continuity (no ties). Define the index sets*

$$P(l) = \{i : D_i(l) > 0\}, \quad N(l) = \{i : D_i(l) < 0\}, \quad p(l) = |P(l)|.$$

Then $T^+(l)$ admits the decomposition

$$T^+(l) = \frac{p(l)(p(l)+1)}{2} + \sum_{i \in P(l)} \sum_{j \in N(l)} \mathbf{1}\{|D_i(l)| > |D_j(l)|\}.$$

Proof. Fix l and suppress (l) in the notation for readability. Let

$$P = \{i : D_i > 0\}, \quad N = \{i : D_i < 0\}, \quad p = |P|.$$

By definition, $T^+ = \sum_{i \in P} R_i$, where R_i is the rank of $|D_i|$ among the nonzero absolute values $\{|D_1|, \dots, |D_b|\}$, ranked from smallest to largest, with no ties by assumption.

For any $i \in P$, the rank R_i can be written as

$$R_i = 1 + \#\{u \in P \setminus \{i\} : |D_u| < |D_i|\} + \#\{v \in N : |D_v| < |D_i|\}. \quad (1)$$

This is the standard “rank equals one plus the number of strictly smaller elements” identity, applied to the multiset of nonzero absolute values.

Summing (1) over $i \in P$ yields three contributions:

$$\sum_{i \in P} R_i = \underbrace{\sum_{i \in P} 1}_{(I)} + \underbrace{\sum_{i \in P} \#\{u \in P \setminus \{i\} : |D_u| < |D_i|\}}_{(II)} + \underbrace{\sum_{i \in P} \#\{v \in N : |D_v| < |D_i|\}}_{(III)}.$$

Term (I): $\sum_{i \in P} 1 = |P| = p$.

Term (II): Among the p positives, each unordered pair $\{i, u\}$ contributes exactly one to the sum because there are no ties. Hence the contribution is $\binom{p}{2}$.

Term (III): By definition,

$$\sum_{i \in P} \#\{v \in N : |D_v| < |D_i|\} = \sum_{i \in P} \sum_{j \in N} \mathbf{1}\{|D_i| > |D_j|\}.$$

Putting these together,

$$T^+ = p + \binom{p}{2} + \sum_{i \in P} \sum_{j \in N} \mathbf{1}\{|D_i| > |D_j|\} = \frac{p(p+1)}{2} + \sum_{i \in P} \sum_{j \in N} \mathbf{1}\{|D_i| > |D_j|\},$$

which is the claimed identity after restoring the explicit dependence on l . □

A.1 Proof of Proposition 5.2 (Power is decreasing in l for Wilcoxon signed-rank)

Let $P(l) = \{i : D_i(l) > 0\}$ and $N(l) = \{i : D_i(l) < 0\}$, with $p(l) = |P(l)|$. Define

$$T^+(l) = \sum_{i=1}^b R_i(l) \mathbf{1}\{D_i(l) > 0\},$$

where $R_i(l)$ is the rank of $|D_i(l)|$ among the nonzero values $\{|D_1(l)|, \dots, |D_b(l)|\}$.

Proposition A.2 (No cross-over $\Rightarrow T^+$ cannot increase). *If $P(l+1) = P(l)$ and $N(l+1) = N(l)$, then $T^+(l+1) \leq T^+(l)$.*

Proof. From Lemma A.1,

$$T^+(l) = \frac{p(l)(p(l)+1)}{2} + \sum_{i \in P(l)} \sum_{j \in N(l)} \mathbf{1}\{|D_i(l)| > |D_j(l)|\}.$$

Case 1: Baseline term. Since $P(l+1) = P(l)$, we have $p(l+1) = p(l)$, so the baseline term $\frac{p(l)(p(l)+1)}{2}$ is identical at l and $l+1$.

Case 2: Positive Negative (PN) comparisons. For $i \in P(l)$, Assumption (A1) implies $|D_i(l+1)| \leq |D_i(l)|$. For $j \in N(l)$, Assumption (A1) implies $|D_j(l+1)| \geq |D_j(l)|$. Hence for each $(i, j) \in P(l) \times N(l)$,

$$\mathbf{1}\{|D_i(l+1)| > |D_j(l+1)|\} \leq \mathbf{1}\{|D_i(l)| > |D_j(l)|\}.$$

Thus every PN indicator is non-increasing.

Conclusion. The baseline is unchanged and the PN sum cannot increase, so $T^+(l+1) \leq T^+(l)$. $\square$

Proposition A.3 (Any cross-over $\Rightarrow T^+$ strictly decreases). *If $P(l+1) \subsetneq P(l)$ (i.e., at least one index that is positive at level l becomes non-positive at level $l+1$), then $T^+(l+1) < T^+(l)$.*

Proof. Let $t \in P(l)$ with $D_t(l+1) \leq 0$. Set

$$p := p(l) = |P(l)|, \quad c_t := \sum_{j \in N(l)} \mathbf{1}\{|D_t(l)| > |D_j(l)|\}.$$

Case 1: Baseline term. When t leaves the positive set, the number of positives drops from p to $p-1$. The baseline term decreases by

$$\frac{(p-1)p}{2} - \frac{p(p+1)}{2} = -p.$$

Case 2: Loss of t 's PN wins. In $T^+(l)$, t contributes c_t units by beating c_t negatives. These are lost at $l+1$.

Case 3: Possible recovery against t . At $l+1$, t is in $N(l+1)$ and can be beaten by remaining positives $i \in P(l) \setminus \{t\}$. Each such comparison contributes at most one unit, so the total possible recovery is at most $(p-1)$.

Combine. Therefore,

$$T^+(l+1) - T^+(l) \leq (-p) + (p-1) - c_t = -1 - c_t \leq -1.$$

Thus a single cross-over strictly decreases T^+ by at least 1. Multiple cross-overs decrease it further. $\square$

Corollary A.4 (Power monotonicity). *Under (A1)–(A2), the WSR statistic satisfies $T^+(l+1) \leq T^+(l)$ point-wise, with strict inequality when a crossover occurs. Since the null distribution of the Wilcoxon signed rank statistic depends only on the number of effective pairs b (and not on l), the critical cutoff c_α is fixed across l , so the monotonicity of $T^+(l)$ directly implies monotonicity of the test power $\beta_j(l)$. Hence the test power*

$$\beta_j(l) = \Pr(T^+(l) \geq c_\alpha \mid H_{1j})$$

is a non-increasing function of l .

A.2 Proof of Proposition 5.1 (Type I Error Decreases in l)

Proposition A.5 (Type I error decreases with l). *Type I error of the WSR test,*

$$\Pr_{H_0}(T^+(l) \geq c_\alpha(b)),$$

is a non-increasing function of l , and strictly decreasing whenever cross-overs $D_i(l) > 0 \geq D_i(l+1)$ occur with positive probability.

Proof. By (A1), $D_i(l+1) \leq D_i(l)$ for all i . Writing $P(l) = \{i : D_i(l) > 0\}$, $N(l) = \{i : D_i(l) < 0\}$, and $p(l) = |P(l)|$, the Wilcoxon decomposition (Lemma A.1) gives

$$T^+(l) = \frac{p(l)(p(l)+1)}{2} + \sum_{i \in P(l)} \sum_{j \in N(l)} \mathbf{1}\{|D_i(l)| > |D_j(l)|\}.$$

If $P(l+1) = P(l)$ and $N(l+1) = N(l)$, then by (A1) every indicator in the double sum is non-increasing, while the baseline term $\frac{p(l)(p(l)+1)}{2}$ is unchanged, hence $T^+(l+1) \leq T^+(l)$ (Proposition A.2). If some $t \in P(l)$ crosses to non-positive at $l+1$, then $T^+(l+1) < T^+(l)$ (Proposition A.3). Therefore,

$$T^+(l+1) \leq_{\text{a.s.}} T^+(l),$$

with strict decrease on any realization exhibiting a cross-over.

Since the Wilcoxon critical value $c_\alpha(b)$ depends only on the number of effective nonzero pairs b (and not on l), for any fixed c we have

$$\Pr_{H_0}(T^+(l+1) \geq c) \leq \Pr_{H_0}(T^+(l) \geq c).$$

Taking $c = c_\alpha(b)$ yields

$$\Pr_{H_0}(T^+(l+1) \geq c_\alpha(b)) \leq \Pr_{H_0}(T^+(l) \geq c_\alpha(b)),$$

which proves that the Type I error is non-increasing in l . Strict inequality holds whenever cross-overs occur with positive probability. $\square$